

The Student Research Workflow

The Student Research Workflow

How to Find, Read, and Organize Research,
and Write Your First Literature Review

Gaurav Tiwari

gauravtiwari.org

Copyright © 2026 Gaurav Tiwari.

All rights reserved. No part of this book may be reproduced in any form without written permission from the author, except for brief quotations in reviews.

First edition, 2026.

gauravtiwari.org

*To everyone with forty-one open tabs
and no idea which one mattered.*

Contents

Preface	ix
I The Landscape	1
1 Research Is a Conversation	2
1.1 The cast of documents	2
1.2 The life cycle of a claim	3
1.3 Sizes of the conversation	4
1.4 What this means for how you work	4
1.5 Three rooms, briefly	4
2 The Anatomy of a Paper	6
2.1 IMRaD, the load-bearing skeleton	6
2.2 The abstract is a contract	7
2.3 Field dialects	7
2.4 Reading order is not page order	7
2.5 A worked walk-through	8
2.6 Papers are versioned too	8
3 Quality and Trust	10
3.1 What peer review actually does	10
3.2 Venue signals, used and abused	10
3.3 Reading the frontier: preprints and gray literature	11
3.4 The trust checklist	11
3.5 A note on your own fallibility	12
3.6 The retraction check, in practice	12
4 From Topic to Question	14
4.1 What makes a question workable	14
4.2 The narrowing funnel	15
4.3 Question types, and what each demands	15
4.4 The pilot search: testing the question	16
4.5 Commitment, in writing, with an exit clause	16
II Finding	18
5 Search Like a Researcher	19
5.1 From topic to questions to queries	19
5.2 Google Scholar, properly	20

5.3	The databases you already pay for	20
5.4	Recall, precision, and which mistake to prefer	20
5.5	A worked search session	21
5.6	When search goes wrong	21
6	The Citation Web	23
6.1	Snowballing backward and forward	23
6.2	Reviews and theses as prebuilt maps	24
6.3	Following people and venues	24
6.4	Citation-graph tools	24
6.5	The web session, worked	25
6.6	Managing the snowball	25
7	Getting the Paper	27
7.1	Why the wall is there	27
7.2	The access ladder	27
7.3	Books, chapters, and theses	28
7.4	The filing moment	28
7.5	The awkward cases	29
8	Knowing When to Stop	30
8.1	Saturation: the signal you can feel	30
8.2	Scope is a sentence you write down	30
8.3	The search log	31
8.4	Time-boxing and the returns curve	31
8.5	The permission paragraph	32
8.6	Two case studies in stopping	32
III	Reading	33
9	The Three-Pass Read	34
9.1	Pass 1: the ten-minute triage	34
9.2	Pass 2: the comprehension read	34
9.3	Pass 3: the adversarial read	35
9.4	Bailing, rereading, and other permissions	35
9.5	The passes meet the pile	36
9.6	Pass 1, in real time	36
10	Reading Critically	38
10.1	Claims versus evidence: the central move	38
10.2	Interrogating methods without a methods degree	38
10.3	A working reader's statistics	39
10.4	Reading against the grain, fairly	39
10.5	The critical read in your notes	40
10.6	Evidence in other dialects	40
11	Notes That Survive	42
11.1	Why highlighting fails, and what works	42
11.2	The paper note template	42
11.3	Scaling the effort to the pass	43

11.4	Notes are for retrieval, so write for search	43
11.5	One worked note	44
11.6	Notes on non-article sources	44
12	The Reading Queue	46
12.1	The queue is tags, not a place	46
12.2	The weekly rhythm	46
12.3	Staying current without drowning	47
12.4	Backlog mercy	47
12.5	The whole pipeline, running	48
12.6	Rhythms for different seasons	48
IV	Organizing	49
13	Your Reference Manager	50
13.1	Zotero, and the one-page honest alternatives	50
13.2	Setup, in one sitting	50
13.3	Metadata hygiene	51
13.4	Organizing: collections, tags, and search	51
13.5	Groups, sharing, and the lab	52
13.6	Migration and maintenance	52
13.7	Beyond papers: what else lives in the library	52
14	The Personal Research Database	54
14.1	From notes to sight	54
14.2	Building the synthesis matrix	54
14.3	Reading the matrix	55
14.4	Keeping it living, keeping it light	55
14.5	Beyond the matrix: when you need more	56
14.6	Matrix pathologies	56
15	Citing Without Tears	58
15.1	What a citation claims	58
15.2	Quotation, paraphrase, and the patchwriting trap	58
15.3	Styles, and what they're actually doing	59
15.4	Citation craft in the sentence	59
15.5	The pre-submission citation pass	60
15.6	A short gallery of citation errors	60
V	The Literature Review	62
16	What a Literature Review Is (and Isn't)	63
16.1	The book report, dissected	63
16.2	The actual specification	63
16.3	The jobs a review does	64
16.4	The genre's variants	64
16.5	Permission to have a voice	65
16.6	Standing on existing reviews, without being buried	65

17 From Notes to Structure	67
17.1 Move 1: choose the organizing logic	67
17.2 Move 2: draft the section map from the matrix	67
17.3 Move 3: allocate the evidence	68
17.4 Move 4: sequence inside sections	68
17.5 Move 5: the field map (one optional page)	69
17.6 A word on the blank-page terror	69
17.7 The advisor in the loop	69
18 Drafting, Revising, and Keeping It Alive	71
18.1 Drafting from the skeleton	71
18.2 The synthesis polish	72
18.3 The honesty audit	72
18.4 The living review	72
18.5 Session logistics	73
19 Review Makeovers	74
19.1 Makeover 1: the book report	74
19.2 Makeover 2: the citation carpet	75
19.3 Makeover 3: the fake gap	75
19.4 Makeover 4: the hedge fog	76
19.5 The pattern across all four	77
19.6 The whole workflow, in one page	77
VI Toolkit	78
A The Paper Note Template	79
B The Synthesis Matrix, Worked	81
C A Field Guide to Sources and Databases	83
D The Literature Review Checklist	85
E Tools, Honestly	87
F Further Reading	89
G The Emergency Protocol	90

Preface

Nobody teaches this. Your degree assumes you can find the right papers, judge them, read them, remember them, and weave them into a literature review, and at no point does anyone show you how. The skills are simply expected to condense out of the air sometime between orientation and your thesis proposal. What actually condenses is the fog: forty open tabs, a downloads folder full of PDFs named `fulltext(3).pdf`, highlights you'll never reread, and a growing dread of the phrase "situate your work in the literature."

I wrote this book because the fog is optional. Research has a workflow the way cooking has a kitchen: find, read, organize, write, each stage feeding the next, each with tools and habits that professionals use daily and students are left to reinvent badly. None of it is hard. All of it is invisible until someone shows you, and then it looks obvious forever.

This book walks the pipeline once, completely, in order. Part I gives you the lay of the land: what the literature actually is, how a paper is built, and how to judge what deserves your trust. Part II is finding: real search technique, the citation web, getting past paywalls legally, and, the skill nobody mentions, knowing when to stop. Part III is reading: the three-pass method, critical reading that goes beyond nodding, and notes that are still useful in eight months. Part IV is organizing: a reference manager doing the remembering, a personal research database, and citation mechanics that never cost you marks. Part V assembles it all into the document the whole pipeline exists for: a literature review organized by ideas, with an honest gap and your voice in charge.

You can read it straight through in a weekend, and the parts also stand alone: if the review is due and the reading is done, start at Part V; if you're drowning in PDFs today, Part IV bails the boat. Every chapter ends with *Put it to work*, a handful of actions you can take the same day, because a workflow you haven't started isn't a workflow, it's a wish.

Two companions from the same shelf, both optional: *How to Write Mathematics* teaches the writing craft for quantitative work, and *LaTeX for Theses and Dissertations* builds the document that your research will eventually live in. This book is about what happens before, and underneath, both.

The literature is not an exam you might fail. It's a conversation that's been waiting for you, and by the last page, you'll know how to walk in, catch up, and say something.

Gaurav Tiwari
gauravtiwari.org

Part I

The Landscape

Research Is a Conversation

Here's the mental model that changes everything, and I want it in place before we touch a single search box or PDF. **The academic literature is not a library. It's a conversation.** A very long, very slow, unusually well-documented argument among thousands of people, running for decades, about what's true and how we know.

Why does the model matter? Because every skill in this book falls out of it. If the literature were a library, your job would be retrieval: find the book with the answer. Students who work from the library model search for "the paper that says X," read it as settled fact, quote it, and wonder why their professor writes "engage with the literature" in the margin. But in a conversation, your job is different. You need to know who's talking, what the current disagreements are, which claims are settled and which are live, and, eventually, where your own small contribution fits. That's what "situate your work" actually means: catch up on the conversation before you speak.

The rest of this chapter gives you the map: what kinds of documents carry the conversation, how a single claim travels through them, and what all of this means for how you'll work.

1.1 The cast of documents

The conversation happens in genres, and each genre has its own speed, its own reliability, and its own use to you. Learning to recognize them on sight is the first practical skill of research, because you read a review article completely differently from a conference paper, and both differently from a textbook.

The journal article is the standard unit of the conversation in most fields: a claim, its evidence, and its methods, written by the people who did the work, checked by peer reviewers, published in a periodical. Articles are slow (months to years from finished work to print) and correspondingly polished. When your professor says "sources," this is mostly what they mean.

The conference paper is the standard unit in computer science and much of engineering, where the field moves too fast for journal timelines. Top conferences there are as selective as top journals, sometimes more. In other fields (much of social science and the humanities), conference papers are early sketches of work that will mature into articles later, and are cited accordingly. The same words, "conference paper," mean different things across the campus, which is exactly the kind of local knowledge this chapter exists to hand you.

The review article is the conversation summarizing itself: an author (usually senior) surveys everything said recently on a topic, organizes it, and points at open questions. Reviews are your best friends for the next several chapters. A good recent review is a map of

the territory drawn by a local, and finding one early can save you twenty hours of wandering. Some fields formalize this as the *systematic review* (medicine especially), with an explicit, reproducible search protocol.

The preprint is a paper shared publicly before (or during) peer review, on servers like arXiv, bioRxiv, SSRN, or a university repository. Preprints are how fast-moving fields actually communicate: physics and machine learning essentially run on arXiv. They carry a specific trade: current but unrefereed. You'll read them, often; Chapter 3 covers how to weigh them.

The monograph, a scholarly book on one topic, is the humanities' flagship genre and appears across fields for mature syntheses. Deep and integrative, but the slowest genre of all. **The textbook** is slower still: by the time a claim reaches a textbook, it's usually a decade past live debate. Fine for learning, almost never the thing to cite for a research claim.

The thesis (someone else's) is an underrated genre: book-length treatment, exhaustive literature review already assembled, methods described at tutorial length because a committee demanded it. Recent theses from good groups are gold mines, especially their bibliography chapters.

Finally, **gray literature**: technical reports, government and NGO publications, working papers, standards documents, datasets, serious blog posts by researchers. Unrefereed but often authoritative in applied topics (try writing about pandemic policy without WHO reports). It has a real place, handled with the care Chapter 3 describes.

Genre tells you how to read before you've read a word. A review orients you, an article argues a claim, a preprint reports from the frontier, a textbook teaches the settled past. Misread the genre and you'll misjudge everything inside it.

1.2 The life cycle of a claim

To see how the genres fit together, follow one claim through the system. Suppose a research group finds evidence that a new teaching method improves exam performance.

First, the finding exists as *talk*: lab meetings, a conference talk, maybe a poster. Fast, invisible to you, mostly. Next, often, a *preprint*: the paper as the authors wrote it, public, free, timestamped, unreviewed. The conversation can now respond, and sometimes does, loudly, before any journal touches it.

Then *peer review*: the authors submit to a journal, an editor recruits two or three researchers in the area, and those reviewers read critically and demand changes (more on what this does and doesn't guarantee in Chapter 3). Months later, the *published article* appears: the version of record, the thing you'll usually cite.

Now the conversation works on the claim. Other groups *cite* it: some build on it, some fail to replicate it, some argue the effect is smaller or conditional. This is the claim's real test, and it takes years. If it survives, it starts appearing in *review articles* as part of "what we know"; a few years later still, it calcifies into *textbooks* and stops being a research topic at all. If it doesn't survive, the record stays messy: the original article still exists, still turns up in searches, still gets cited by people who missed the follow-ups. **Publication is a claim's birth certificate, not its verdict.** The verdict is distributed across the citing literature, which is why Chapter 6 teaches you to check what happened *after* a paper, not just what it says.

Notice, too, what this life cycle implies about dates. A 2009 paper isn't "too old" and a 2026 preprint isn't "better because it's newer." Each sits at a different point in the cycle: the old paper has a citation record you can interrogate; the new preprint is the frontier, unvetted.

A good literature review, we'll see in Part V, deliberately spans the cycle: foundations, current state, frontier.

1.3 Sizes of the conversation

One more piece of orientation, because it explains feelings you'll have in Part II. Conversations come in sizes. Some topics are enormous rooms: tens of thousands of papers, dozens of reviews, whole journals dedicated to them. Others are a corner table: forty papers, three groups, everyone citing everyone. Your project lives somewhere specific in this range, and misjudging the size produces the two classic failure feelings. In a big conversation, the feeling is drowning: every search returns thousands of hits, and the fix is narrowing (Chapter 5) and leaning on reviews (Chapter 6). In a small one, the feeling is starving: searches return almost nothing, you conclude you're either a genius or lost, and the fix is widening the net to adjacent conversations, because a topic with genuinely no literature is rare, while a topic whose literature speaks a different vocabulary than you searched is extremely common.

Part of your first week's work on any project is simply sizing the room, and Chapter 8 turns that from a feeling into a procedure.

1.4 What this means for how you work

The conversation model sets up everything that follows, so let me draw the consequences explicitly, as commitments this book will keep making concrete:

1. **You'll read for positions, not just facts.** Every paper is a move in an argument: it agrees, disagrees, extends, qualifies. Your notes (Chapter 11) will record what a paper *does* to the conversation, not just what it says, and your literature review (Part V) will be organized around exactly those moves.
2. **You'll care about relationships between papers** as much as papers themselves. Who cites whom, who answered whom, which paper started the thread: the citation web (Chapter 6) is the conversation's own index of itself.
3. **You'll evaluate speakers, not worship them.** Peer review is a filter, not a guarantee; venues have reputations; claims have afterlives. Chapter 3 makes you a fair but unintimidated judge.
4. **You'll keep records like someone planning to speak.** Because you are. The reference manager, the notes, the synthesis matrix (Part IV): all of it exists so that when your turn comes, in the review, the thesis, the defense, you can cite the conversation accurately and join it with confidence.

There's one more consequence, and it's the emotional one. Students walk into the literature feeling like trespassers in an experts-only room. The conversation model reframes it: nobody in that room knows everything either; every one of them once stood where you stand; and the genres exist precisely so newcomers can catch up. Reviews are the room's own welcome desk. Theses are its guided tours. You're not trespassing. You're the newest participant, and the whole apparatus, built over three centuries, is set up to brief you.

1.5 Three rooms, briefly

Because this book's running example lives in social science, here's the conversation glimpsed in three other rooms, so you can recognize yours.

Machine learning runs at preprint speed: the unit is the conference paper, posted to arXiv months before the conference, discussed publicly within days. Reviews lag the frontier badly, so orientation comes from survey papers and curated reading lists, and the citation web (Chapter 6) matters more than any database. A “recent” paper is from last spring.

Clinical psychology runs on journals and formal synthesis: randomized trials accumulate slowly, systematic reviews and meta-analyses adjudicate, and the field’s replication reckoning means methods sections get read with real suspicion (Chapter 10 will show you how). A “recent” paper is from the last five years, and a 1995 measurement-validation study may still be load-bearing.

History runs on monographs and archives: the article matters, but the book is the career unit, book reviews are serious critical literature, and “evidence” means sources read in context. A “recent” intervention may be a decade old, and the conversation’s fault lines run through interpretation, not statistics.

Same structure everywhere: genres, speeds, fault lines, and a newcomer’s apparatus for catching up. Only the constants change.

Put it to work

- Write your project topic at the top of a page. Under it, list what you already suspect are the “rooms”: the two or three conversations it touches. Guesses are fine; Part II will correct them.
- Find one **review article** adjacent to your topic (search your topic plus the word “review” or “survey”; formal technique arrives in Chapter 5). Skim only its section headings and its reference count. That’s the shape of your room.
- Identify which genre dominates *your* field’s conversation: ask an advisor or senior student, “where does the real work appear here, journals or conferences?” The answer calibrates everything in Parts II and III.
- Pick one recent **thesis** from your department on anything nearby (departments list them; so do repositories). Note how its literature-review chapter is organized. You’re looking at your own destination.

2

The Anatomy of a Paper

A research paper looks like an essay that ate its vegetables: sections, subsections, an abstract, forty references. It isn't an essay at all. It's a structured document with fixed compartments, and every compartment answers a specific question a skeptical reader is entitled to ask. Once you know which question lives where, you can open any paper in any field and navigate it in seconds, which is the foundation for everything Part III will teach about reading. This chapter is the floor plan.

2.1 IMRaD, the load-bearing skeleton

Most empirical papers, across the sciences and social sciences, follow IMRaD: **I**ntroduction, **M**ethods, **R**esults, and **D**iscussion. The structure isn't bureaucratic tradition; each section is the answer to one question:

- **Introduction:** *why should anyone care, and what's the question?* Context, the gap in the conversation (Chapter 1's language is literal here; papers say "however, no prior work has..."), and usually a statement of the contribution.
- **Methods:** *what exactly did you do?* In enough detail, in principle, for someone to redo it. Study design, participants or materials, procedures, analysis plan.
- **Results:** *what happened?* The findings, ideally reported without spin: numbers, tables, figures, statistics.
- **Discussion:** *what does it mean?* Interpretation, connection back to the literature, limitations, implications.

The deep insight, the one that changes how you read forever, is that **the sections carry different kinds of statement**. Results are (supposed to be) observations. Discussions are arguments. Introductions are framing. The same paper can have solid results and an overreaching discussion, and sophisticated readers routinely accept a paper's data while rejecting its interpretation. Chapter 10 builds on exactly this: you'll learn to locate a claim's compartment before deciding how much to trust it.

Around the skeleton, the standard organs: the **title** (a claim or a topic, compressed), the **abstract** (next section), **keywords**, the **author list** and affiliations (more informative than it looks: labs have styles and track records), the **references**, and increasingly a **data or code availability statement**, plus **supplementary material**, the overflow basement where extra figures, robustness checks, and gory method details live. Modern reading often requires a trip to the basement, and it's where skeptical readers check whether a slick figure survives contact with the fuller data.

2.2 The abstract is a contract

The abstract is the paper in miniature: one paragraph, typically 150 to 300 words, almost always built as micro-IMRaD: a sentence or two of context, the question, the approach, the key result (often with the actual number), and the takeaway.

Learn to read abstracts *structurally*. Here's a made-up but typical one, dissected:

“Spaced retrieval practice improves long-term retention, but its interaction with sleep remains unclear [context and gap]. We randomized 240 undergraduates to spaced or massed practice, with testing before or after a sleep interval [method]. Spaced practice improved 30-day retention by 21% (95% CI: 14–28%), with no significant sleep interaction [results]. Retrieval spacing, rather than sleep-dependent consolidation, appears to drive the benefit [interpretation].”

Four moves, four sentences. Once you see the template, you can extract a paper's entire skeleton from its abstract in under a minute, which is precisely what pass 1 of the three-pass read (Chapter 9) will ask you to do at volume.

Why call it a contract? Because the abstract is the authors' public promise of what's inside, and one of your recurring critical-reading jobs is checking whether the paper keeps it. Abstracts are the most-read, most-optimized, and therefore most-inflated part of any paper: effects grow adjectives, caveats evaporate, “suggests” becomes “shows.” The gap between abstract and results section, when you find one, is diagnostic, and Chapter 10 will send you hunting for it deliberately.

2.3 Field dialects

IMRaD is the majority dialect, not the law. You'll meet at least three other body plans, and recognizing them prevents the classic disorientation of “where are the Methods?”

Mathematics and theoretical work (Chapter 1's conversation running on proofs): introduction, definitions and preliminaries, theorem-proof-theorem, discussion. The “results” are the theorems; the “methods” are the proofs; the evidence is logical rather than empirical. My *How to Write Mathematics* lives entirely inside this dialect.

Computer science and engineering: often IMRaD-with-renovations: an explicit *Related Work* section (the introduction's gap-argument, promoted to its own room), a *System* or *Approach* section instead of Methods, an *Evaluation* instead of Results, and frequently a *Limitations* or *Threats to Validity* section that's required reading.

Humanities: continuous argument, few or no section headings, evidence woven as quotation and close reading, the thesis often arriving after considerable throat-clearing. Navigation here is by paragraph function rather than headings: opening moves, counterpositions, the author's claim, support, response to objections. The compartments exist; they're just unlabeled.

When you cross into an unfamiliar field, and Part II's searching will absolutely take you across, spend your first paper identifying the local body plan before judging content. A paper isn't disorganized because its organs are in different places than your home field keeps them.

2.4 Reading order is not page order

Now the payoff. Because papers are compartmentalized, nobody who reads for a living reads them front to back. Different missions route differently through the same document, and here are the three routes you'll use constantly:

The triage route (is this paper relevant to me?): title, abstract, then straight to the figures and tables, then the introduction's final paragraph, where most papers state their contribution in list form. Total: three to five minutes, and usually decisive. This is pass 1 of Chapter 9, and you'll run it dozens of times per project.

The claim-checking route (someone cited this for claim X; does it really say that?): abstract, then search the Results and Discussion for the claim, then check the method just enough to see what kind of evidence backs it. You'll run this route every time you follow a citation from another paper, and you'll be shocked, genuinely, at how often the cited paper says something narrower, more conditional, or simply different from what the citing paper implied. That discovery is a rite of passage, and Chapter 15 will make "cite only what you've verified" one of your professional commitments.

The comprehension route (I need to actually understand this): figures first, seriously. In most empirical papers the figures and their captions carry the core findings, and authors design them to. Then Results with the figures beside you, then Methods with your skeptic's hat on, then Introduction and Discussion last, *after* you've formed your own view of the data, so the authors' framing lands on a mind that's already thinking rather than one that's still blank. This inversion, data before framing, is the single most useful reading habit I know, and Chapter 10 is built on it.

Page order is the authors' presentation of their work. Reading order is your interrogation of it. The moment you stop reading papers like novels, they stop being exhausting.

2.5 A worked walk-through

Take fifteen minutes and try the comprehension route on one real paper near your topic (any will do; the review article you found in Chapter 1 can supply a promising reference). Here's the walk, step by step:

1. **Title and abstract** (2 minutes): dissect the abstract into its four moves, on paper. What's the promised contribution, in one sentence you write yourself?
2. **Figures and tables** (4 minutes): for each, read the caption, then ask: what comparison is this showing? What's the axis, the unit, the size of the effect? Can you state the main finding from the figures alone?
3. **Results** (4 minutes): find the paragraph behind the most important figure. Does the text's claim match what the figure shows, in size and confidence?
4. **Methods** (3 minutes, skim): who or what was studied, how many, measured how? Anything that surprises you about how the evidence was actually produced?
5. **Introduction and Discussion** (2 minutes): now read their framing. Where do they place the work in the conversation? Do the limitations they admit match the ones you noticed?

Fifteen minutes, and notice what you have: a structural summary, an independent read of the evidence, and a first critical reaction, before you've "read the paper" in the naive sense at all. That experience, scaled and systematized, is all of Part III.

2.6 Papers are versioned too

One anatomy fact that surprises students: a paper is often several documents wearing one title. The preprint (v1), its revisions (v2, v3, sometimes years apart), the accepted manuscript, and the version of record can differ in results, in framing, and occasionally in conclusions;

peer review exists to force changes, and it does. Practical consequences: check which version you're holding (repository PDFs announce it; arXiv numbers versions explicitly), prefer the version of record once it exists (Chapter 7), and when a claim matters, glance at whether it survived into the published version, because citing a v1 claim that v3 walked back is a small, avoidable embarrassment with your name on it. Your reference manager records the version (Chapter 13), and your notes record which one you read (Chapter 11). Fields differ in how much this matters: in preprint-native fields, versions are daily weather; in journal-first fields, the published PDF is usually the only one you'll meet. Either way, "the paper" is an anatomy lesson of its own: not one immutable object, but the current state of a document that argued its way through review.

Put it to work

- Run the fifteen-minute walk-through above on one paper today. Keep your notes; Chapter 11 will show you what they should grow into.
- Dissect the abstracts of three papers in your field into the four moves. If a move is missing, note which; you're learning your field's dialect.
- Identify your field's body plan (IMRaD, theorem-proof, renovated IMRaD, continuous argument) and its non-negotiable extra organs (Related Work? Threats to Validity? Data availability?).
- Bookmark one paper's supplementary material and skim it once, just to see what basements contain in your field.

3

Quality and Trust

Every source you use is a bet: you're staking a piece of your own credibility on someone else's work. Students handle this bet in two bad ways. Some trust everything with a DOI, treating "published" as "true." Others, after their first exposure to scientific scandal, trust nothing and drown in suspicion. Both are ways of not thinking. This chapter gives you the working researcher's alternative: a calibrated, checkable sense of how much weight each kind of source can bear.

3.1 What peer review actually does

Peer review, honestly described: an editor sends a submitted paper to (usually) two or three unpaid researchers in the same area. They read it, write critiques, and recommend: reject, revise, or accept. The authors revise, sometimes through several rounds. Median time from submission to acceptance runs from months to over a year, field depending.

What this process is good at: catching obvious methodological holes, forcing clarity, demanding missing analyses and caveats, filtering work that's clearly below a venue's bar. A peer-reviewed paper has survived hostile expert reading, which is genuinely worth something, and it's the reason the published article is the conversation's unit of record.

What it is not: verification. **Reviewers almost never rerun experiments, reanalyze raw data, or check computations.** They review the *report*, not the work. Peer review therefore cannot reliably catch subtle statistical problems, honest errors deep in analysis pipelines, or fraud, and empirically it doesn't: the replication crisis (whole literatures of findings, especially in psychology and biomedicine, failing to reproduce at scale) and every fraud scandal you've read about happened *in peer-reviewed venues*. Retractions, the system's formal error-correction (a journal marking a published paper as unreliable), run years behind, and retracted papers keep accumulating citations from readers who never learned.

So calibrate: **peer review is a competence filter, not a truth stamp.** It moves a paper from "someone claims" to "a claim that survived expert scrutiny of its reasoning." The truth question is answered, slowly, by the conversation: replication, extension, citation in reviews, the afterlife you learned to check in Chapter 1. A single unreplicated paper, however prestigious its journal, is one data point, and your literature review should treat it as one.

3.2 Venue signals, used and abused

Where a paper appears carries information, and the information is routinely over-read. The honest hierarchy of what venue tells you:

Field reputation is the real signal. Every field has venues its members respect, and the ranking is common knowledge *inside* the field and invisible outside it. One conversation with an advisor (“which journals here make people sit up?”) transfers it. This beats every metric below.

Impact factor, the famous number, is the mean citations per recent article in a journal. It measures the journal’s average visibility, not your paper’s quality: citation counts within a journal are wildly skewed (a few blockbusters, a long quiet tail), and fields differ so much in citation culture that cross-field comparison is meaningless. Use it, at most, as a coarse tier signal within a field you don’t know yet. The same goes for per-paper citation counts: highly informative after years (hundreds of citations means the conversation engaged), nearly meaningless for anything recent, and always contaminated by citations that disagree, a nuance Chapter 6 exploits productively.

Predatory journals are the abuse case you must be able to spot: outlets that mimic the form of scholarship (editorial boards, “peer review,” DOIs, fees) while accepting essentially anything for money. They exist because careers are scored on publication counts, and they’re common enough that you *will* encounter them in searches. Warning signs, any two of which should make you check further: a title suspiciously close to a famous journal’s; grandiose scope (“International Journal of Science and Advanced Innovation”); solicitation spam; promised review times of days; an “editorial board” of unknown or unconsenting names; no institutional affiliation you can verify. The positive check is easier: is the journal in the major indexes your library subscribes to, does your field’s community publish there, has your advisor heard of it? When in doubt, don’t build on it, and don’t cite it except as an object of study.

3.3 Reading the frontier: preprints and gray literature

Chapter 1 put preprints at the frontier of the conversation; here’s how to weigh them. A preprint deserves the same structural reading as any paper (Chapter 2) plus two extra checks: **what happened next** (many preprints have a published version of record a year later; always look, and cite the version of record when it exists, Chapter 15), and **who’s engaging** (a preprint that respected groups are already citing and building on has informal review; a preprint nobody has touched in three years is a flag, though sometimes just an orphaned good idea). In preprint-native fields, mature researchers read preprints constantly and cite them with dates and version numbers; in conservative fields (medicine above all), preprint claims about anything consequential need the published version or corroboration before they carry weight in your writing.

Gray literature earns trust through **institutional process** instead of peer review: a central bank’s working paper, a standards body’s specification, a WHO technical report each have internal review cultures, and their authority is often exactly what your applied topic needs. The checks: is the issuing body authoritative *for this claim*, is the document current, is it the final version, and is there an underlying peer-reviewed source you should cite alongside or instead? A serious blog post by a known researcher can be worth reading and even citing (fields vary); an anonymous Medium post is not a source, it’s a lead.

3.4 The trust checklist

Here’s the whole chapter as the working protocol you’ll actually run. For any source that might carry weight in your project, four questions, thirty seconds each:

1. **What kind of thing is this?** Genre and venue: refereed article, conference paper (which culture?), preprint, review, gray literature, blog. Sets the baseline (Chapters 1 and this one).
2. **Who says so, on what evidence?** Authors and groups you can verify; and locate the claim's compartment (Chapter 2): is it a result, or discussion-section interpretation?
3. **What happened next?** Citations, replications, follow-ups, corrections, retraction check for anything load-bearing (retraction databases exist and take seconds). Chapter 6 makes this mechanical.
4. **Does anything conflict?** Funding and conflicts of interest statements (standard sections now); and does the claim sit oddly with the rest of the literature you're assembling? An outlier isn't wrong, but it's an outlier, and your notes (Chapter 11) should say so.

Two calibration rules complete the protocol. **Weight scales with load.** A passing mention in your review can rest on lighter sourcing than the finding your whole project builds on; for anything load-bearing, insist on the strongest available genre plus a healthy afterlife. And **prefer the conversation to the soloist:** one review plus its strongest three primary sources beats ten scattered singletons, always.

Trust is not binary and not free. Every source gets a genre, a context, and an afterlife check, and earns exactly the weight those support. "Peer-reviewed" opens the door; it doesn't set the table.

3.5 A note on your own fallibility

One last calibration, inward this time. The failure mode that actually damages student projects isn't citing a predatory journal; it's **confirmation shopping:** searching until you find the paper that says what your argument needs, trusting it because it agrees with you, and stopping. Every technique in this chapter can be corrupted into a weapon for dismissing inconvenient evidence and a rubber stamp for convenient evidence. The antidote is procedural, and Part V builds it in: your synthesis matrix (Chapter 14) will record disagreeing sources side by side, and your literature review (Chapter 17) will be graded, by every good reader, precisely on how honestly it handles the papers that *don't* support you. Start practicing now: for every source you're delighted to find, ask what the strongest paper on the other side says, and go find it.

3.6 The retraction check, in practice

Chapter 1 said retractions lag and retracted papers keep collecting citations; here's the ninety-second protocol that keeps them out of your review. For any load-bearing source: search the Retraction Watch database (free, searchable by author and title), glance at the paper's page on the publisher site (retractions and corrections are posted there, and DOIs resolve to the notice), and let your reference manager help (plugins exist that flag retracted DOIs in your library; ask Appendix E's criteria of any you adopt). Expressions of concern, the milder cousin, deserve a note in your LIMITS field and a weight reduction, not automatic exile.

If a cornerstone of your project does turn out retracted or seriously corrected mid-project, the response is scholarly, not panicked: find out *why* (fraud differs from honest error differs from a statistical dispute), check whether the finding survives elsewhere (the conversation usually re-tested it: Chapter 6's forward snowball tells you), and say what

happened in your review if the episode is itself informative. A field correcting itself is not an embarrassment to cite; it's the system working, and examiners respect a student who handled it visibly.

Put it to work

- Ask your advisor or a senior student for the top three to five respected venues in your area. Write them down; this list recalibrates every search result you'll see in Part II.
- Take one load-bearing paper for your project and run the four-question protocol on it, in writing, including the what-happened-next check. Budget ten minutes.
- Find the funding and conflict-of-interest statements in two papers you've already collected, just to build the reflex of looking.
- Pick one claim you *want* to be true for your project. Spend fifteen minutes deliberately hunting the best evidence against it. Whatever you find goes in your notes, permanently.

4

From Topic to Question

There's a step between "I have a topic" and everything else in this book, and projects that skip it pay for months. A topic is a region: *social media and mental health*, *renewable energy storage*, *the novels of Toni Morrison*. You can't search a region systematically (Chapter 5 needs concepts), can't scope it (Chapter 8 needs boundaries), can't judge sources against it (relevant to *what?*), and can't finish it, ever. A question, by contrast, is an address: specific enough to search, scoped enough to saturate, and, one day, answerable enough to stop. This chapter is the craft of getting from region to address, and since the rest of this book runs on a worked example, this is also the chapter where that example gets built honestly, from its topic up.

4.1 What makes a question workable

Not every question deserves your semester. The five properties that separate workable research questions from essay prompts, each a test you can run in a minute:

1. It's decidable by evidence you can reach. There must exist observations, documents, data, or arguments that would count as an answer, and they must be accessible to *you*, at your level, in your timeframe, with Chapter 7's ladder. "Does social media harm teenagers?" fails twice: *harm* isn't operationalized (what would count?), and the definitive causal evidence doesn't exist for anyone, let alone a student with eight months. "What does the logged-usage literature find about passive use and adolescent anxiety?" passes both.

2. It's scaled to the container. A course paper holds a question the size of one fault line; a master's thesis holds a question plus an empirical contribution; a PhD holds a small family of them. The classic student error is a dissertation-sized question in a term-paper container, and the tell is a scope statement (Chapter 8) that keeps bursting. Scale down by adding constraints, population, period, method, setting, until the container closes.

3. It connects to a live conversation. Your question should land *in* a room from Chapter 1: people are discussing it, or discussing something your question extends. This isn't conformism; it's what makes the literature usable at all (orientation, methods to borrow, results to compare) and what makes your answer matter to someone. The pilot searches below test this cheaply.

4. It's open, genuinely. If every source agrees, your question is a summary assignment, fine for learning, thin for research. The best student questions sit at fault lines (Chapter 14 will find them mechanically), where the honest answer is "the evidence splits, and here's what the split tracks."

5. You can stand it for months. Underrated and real: you'll spend more hours with this question than with most people in your life. Some part of it should genuinely bug you. Curiosity survives week nine; dutiful topic-selection doesn't.

4.2 The narrowing funnel

The move from region to address is a funnel with named stages, and it's worth doing on paper, in one sitting, even badly, because every stage is revisable and only the written version can be revised. Watch the running example descend:

TOPIC	social media and mental health
ASPECT	which harm? -> anxiety (not depression, not self-esteem, not sleep)
WHO/WHERE	adolescents 12-18 (not adults, not children)
WHAT KIND	everyday use (not clinical addiction; not cyberbullying, its own literature)
RELATION	does use relate to anxiety -- and which use? the passive/active distinction keeps surfacing
QUESTION	What does the evidence say about the relationship between social media use, especially passive use, and anxiety in adolescents -- and how much does measurement quality change the answer?

Each stage is a choice among alternatives, and the choices are the question's DNA: change "anxiety" to "sleep" or "adolescents" to "first-year undergraduates" and a different, equally legitimate project exists. Two mechanics make the funnel work:

Narrow with the literature, not against it. The stages after ASPECT came from pilot contact with the conversation: the passive/active distinction, the measurement issue, the cyberbullying-is-separate discovery. You cannot funnel well from an armchair; you funnel, search a little, funnel again. Which is why this chapter sits *before* Part II but its work interleaves with Part II's first sessions.

Keep the discarded branches. The aspects you didn't choose (sleep, self-esteem) become your scope statement's exclusion lines (Chapter 8) and your review's "beyond our scope" sentences (Part V). Nothing in the funnel is wasted; it's pre-written boundary prose.

4.3 Question types, and what each demands

Questions come in types, each with its own evidence requirements and its own downstream workflow emphasis. Typing your question early tells you what you're signing up for:

- **Descriptive:** what's the state of X? (prevalence, patterns, trends). Demands: good measurement literature, representative data. Reviews of descriptive work lean on the measurement columns of your matrix.
- **Comparative:** does X differ across groups, settings, eras? Demands: comparable measures across the compared things, and vigilance about confounds (Chapter 10's "compared to what" cuts both ways).
- **Relational:** does X track Y? The correlational workhorse. Demands: honesty about direction and third variables; your review will live at the causality fault line.
- **Causal:** does X produce Y? The most wanted and least available answer. Demands: experimental or clever quasi-experimental designs, and if the field lacks them, your question may need to *become* relational-plus-honest, which is exactly what our running example did with its "and how much does measurement change the answer" tail.
- **Mechanistic:** how, by what pathway? Demands: process evidence (the qualitative and theory dialects of Chapter 10 carry real weight here).

- **Interpretive** (the humanities' home): what does X mean, how should it be read? Demands: primary sources, engagement with existing readings, and an argument standard rather than a statistical one.
- **Design** (engineering, CS, applied fields): can we build X to do Y within Z? Demands: benchmarks, baselines, and a related-work section that proves the gap is real.

Mixed types are normal; the funnel's RELATION stage is where you commit to the primary one, because it sets the evidence bar your review must clear.

4.4 The pilot search: testing the question

Before committing, spend one hour running the question through two cheap tests, using tools you'll master properly in Part II:

The room test. Search the question's concepts (a proto concept table, Chapter 5). You're hoping for a *mid-sized* room: a few hundred to a few thousand hits, at least one review, recognizable recurring names. A barren room at every vocabulary means property 3 fails (or you've found a genuine frontier, which for a first project is usually a polite way of saying "unsupervised"); an ocean means the funnel needs another stage.

The neighbor test. Find the three closest existing works. If one *is* your question, answered recently and well: excellent news at week one (imagine week twenty), and the funnel runs again from their limitations sections, which are lists of adjacent open questions helpfully published by experts. If the neighbors are near-but-not-on, you've confirmed a live conversation with an open seat: property 3 and 4, both checked.

4.5 Commitment, in writing, with an exit clause

The chapter's output is one page: the question (one sentence), its type, the funnel that produced it (five lines), the pilot results (three lines), and, for supervised work, a sign-off. Show it to your advisor *before* the heavy search weeks of Part II; redirecting a question costs minutes now and weeks later.

Then the exit clause, because questions legitimately evolve: **the question may change when the literature teaches you something, consciously, on paper, once or twice, and not by drift.** Keep a question log next to the search log: dated versions, one line on why each changed ("v2: narrowed to passive use; the literature splits on it and the general question was answered by Rev2020"). Drift, the unconscious version, produces the saddest artifact in student research: a review that answers three half-questions badly. The log is how you notice drift while it's still a decision.

A topic is where you'll look; a question is what you'll answer. Funnel deliberately, type it, pilot it, write it down, and let it change only on purpose. Every tool in the next fourteen chapters points at whatever sentence this chapter produced, which is reason enough to make it a good one.

Put it to work

- Run your topic through the funnel on one page, choosing at every stage and keeping the discarded branches. Time-box: forty minutes, revisable forever.
- Type your question against the seven types and write one sentence on what that type demands of your evidence.

- Do the one-hour pilot: room test and neighbor test, with three lines of findings each. If either fails, funnel again tonight; it's cheap tonight.
- Start the question log with v1, dated, and book the ten-minute advisor sign-off before your first full search week.

Part II

Finding

5

Search Like a Researcher

Watch a student search the literature and you'll usually see one move: type the essay question into Google Scholar, read the first page of results, feel overwhelmed, repeat with slightly different words. Watch a researcher and you'll see something else entirely: a loop of query, scan, harvest, refine, running across two or three tools, leaving a written trail. The difference isn't talent. It's technique, about an hour's worth, and this chapter transfers it.

5.1 From topic to questions to queries

Searches fail before the search box, at the framing step. A topic ("social media and mental health") isn't searchable; it's a continent. The pipeline that works: topic → specific questions → concepts → query vocabulary.

Take the topic above. Specific questions carve it: *Does Instagram use correlate with anxiety in teenagers? Do usage limits help? Is passive scrolling worse than active posting?* Pick one. Now decompose it into concepts, and for each concept, brainstorm the vocabulary the literature might use, because **the literature's words are rarely your words**, and this mismatch is the number-one cause of "there's nothing on my topic":

```
Concept 1: social media  -> "social media", Instagram,
    "social networking sites", SNS
Concept 2: mental health -> anxiety, depression,
    "psychological well-being", "internalizing symptoms"
Concept 3: teenagers      -> adolescen* (adolescent, adolescence),
    youth, "secondary school"
```

That concept table is your search's source code. Queries are just combinations of it: one term per concept, AND between concepts, OR between synonyms, quotes around phrases, * for word-stem wildcards where the database supports them:

```
("social media" OR Instagram) AND (anxiety OR "well-being")
AND adolescen*
```

You'll iterate: every useful paper you find teaches you new vocabulary (the field says "problematic smartphone use"! It says "fear of missing out"!), and new vocabulary goes back into the table and generates better queries. Searching is a loop, not a lookup, and the table is what makes the loop converge instead of wander.

5.2 Google Scholar, properly

Scholar is the default front door for good reasons: it indexes nearly everything (articles, preprints, theses, books), links to PDFs, and shows citation counts inline. Use it deliberately:

- **Quotes and minus** work as you'd hope: "spaced retrieval" forces the phrase, -review drops a term.
- **Cited by** and **Related articles**, under every result, are the citation web's front door; Chapter 6 lives there.
- **Date filters** (left sidebar) matter more than students expect: a "since 2020" pass shows you the live conversation; an all-time pass sorted by citations shows you the foundations. Run both, deliberately, for every question.
- **The profile pages** (authors' own) list everything by a researcher, with counts: the fastest way to survey a key person's work.
- Know its limits: Scholar ranks by a citation-weighted relevance that favors older, famous work; its coverage is broad but its metadata is scruffy; and it has no controlled vocabulary, so it only ever matches *words*, which is exactly why your concept table carries synonyms.

5.3 The databases you already pay for

Your tuition buys access to curated databases that beat Scholar precisely where it's weak, and most students graduate without touching them. The general pattern: field-specific coverage, clean metadata, powerful filters, and, the crown jewel, **controlled vocabulary**: human-assigned subject tags, so you can find every paper *about* a concept regardless of the words its authors chose. Medicine's PubMed has MeSH terms; psychology's PsycINFO has its thesaurus; engineering has IEEE Xplore and Scopus; humanities have MLA International Bibliography and JSTOR; economics has EconLit; and the multidisciplinary giants, Scopus and Web of Science, cover everything with strong citation tooling. (Appendix C maps fields to databases in detail.)

The practical protocol: do your exploratory loops in Scholar, where speed lives; then, once your concept table is mature, run one **systematic pass** in your field's flagship database with proper filters (peer-reviewed only, date range, subject tags). The systematic pass is also where your search becomes *reportable*: thesis methods sections and systematic reviews say things like "we searched PsycINFO with the following terms," and Chapter 8's search log makes that sentence free to write.

The library's other superpower is people. Subject librarians exist, know these databases the way you know your phone, and will sit with you for an hour for nothing. One session in week two of a thesis is worth ten hours of flailing; I have never met a student who regretted booking it, only students who discovered librarians in their final semester and mourned.

5.4 Recall, precision, and which mistake to prefer

Search theory in one paragraph, because it explains your two bad feelings. **Precision** is the fraction of your results that are relevant; **recall** is the fraction of relevant papers you found. They trade off: broad queries flood you (high recall, low precision, the drowning feeling); narrow queries run dry (high precision, low recall, the starving feeling, and its dangerous cousin: missing the paper your examiner will ask about).

The professional habit is to **err toward recall early, precision late**. Early in a project, run broad, skim hard (pass-1 reading, Chapter 9, makes skimming cheap), and let the concept table grow. Late in a project, run narrow, verifying you've covered specific fronts. And through both phases, remember that keyword search is only one of the two great discovery engines: the other, following the citation web, is next chapter, and it's the one that finds what your keywords never will.

You're not looking for "some sources." You're mapping a conversation. Broad first to find its edges, narrow later to fill its gaps, and write down every query, because a search you can't remember is a search you'll redo.

5.5 A worked search session

Here's what forty-five real minutes look like, so the loop is concrete. Project: the Instagram-and-teen-anxiety question above.

Minutes 0–5. Build the concept table (as above). Three concepts, ten terms.

Minutes 5–15. Scholar, broad query, all time, sorted by relevance. Skim two pages of titles. Harvest: two obvious must-reads (into the reference manager, Chapter 13, takes one click each), one review article ("Social media and adolescent mental health: a review"), jackpot, and four new vocabulary items for the table ("screen time," "FOMO," "passive use," "problematic use").

Minutes 15–25. Same query, "since 2022" filter: what's the conversation doing *now*? Two more harvests, including a big recent study everyone seems to cite. Notice a dissenting title ("...little evidence for..."): harvest that especially (Chapter 3's antidote in action). *Minutes 25–35.* Rebuild the best query with the new vocabulary, run once more; diminishing returns already, which is informative (Chapter 8 will formalize why). Log the queries.

Minutes 35–45. Switch to PsycINFO. Find the thesaurus terms for your concepts, run the systematic version, filters on. Different results surface, notably an older foundational study Scholar buried. Harvest, log, stop for the day.

Ten to twelve papers harvested, a vocabulary that doubled, two tools cross-checked, a written trail, and, crucially, a *stopping point*: the session ended on schedule, not in exhaustion. That's a search session. You'll run a handful of them per project, not a hundred.

5.6 When search goes wrong

Three failure patterns cover most bad sessions, each with a mechanical fix:

Near-zero results. Almost never means no literature; almost always means vocabulary mismatch or an over-stacked query. Fixes, in order: drop your rarest AND-concept (three concepts maximum to start); replace your words with the field's (find one on-topic paper by any route, harvest its keywords and title vocabulary); check you're in a database that covers the field at all (Appendix C). If a stripped-down, right-vocabulary query in the right database still returns crickets, congratulations: you may have found an actual gap, and Chapter 8 tells you how to claim it responsibly.

Thousands of results. Add the missing concept (usually population or context), switch the session's goal from papers to *reviews* (add "review OR meta-analysis" and harvest maps instead of bricks), or tighten a term from concept to phrase ("well-being" → "anxiety symptoms"). Drowning is a scoping signal, not a reading assignment.

The homonym trap. "Transformer," "attachment," "resilience," "inflation": when your term means something else in a bigger field, your results fill with the bigger field.

Fix by adding a disambiguating concept from your side (“attachment AND infant”), or excluding theirs (“transformer -electrical”), and note the collision in your search log so future queries inherit the repair.

One more honest note: most databases index English-language venues best, and whole literatures exist in other languages. For most student projects this is a scope line to declare (Chapter 8), not a failure; for some topics (area studies, education policy, local health), the non-English literature is the literature, and your librarian knows the regional databases.

Put it to work

- Build the concept table for your project: concepts as columns, synonym vocabulary under each. Twenty minutes, and it becomes the most-edited page in your notes.
- Run the forty-five-minute session above on your own question, including the systematic pass in one library database. Log every query with its date and result count.
- Book the subject librarian. Actually book them; the button is on the library site right now.
- Run one all-time, citation-sorted Scholar pass and one since-last-two-years pass on your main query, and note the three foundations and three frontier papers they respectively surface.

The Citation Web

Keyword search finds papers that share your vocabulary. But the best papers for your project might use vocabulary you've never heard, sit in journals you'd never check, or predate your search window entirely. They're findable anyway, because the literature carries its own index: every paper's reference list points backward to what it built on, and every citation of it points forward to what built on it. Learning to walk this web is the second great discovery engine, and for depth (as opposed to breadth), it beats keyword search every time.

6.1 Snowballing backward and forward

The technique has an unglamorous name, snowballing, and two directions.

Backward snowballing is mining reference lists. Take your best paper so far, the one closest to your project, and read its bibliography as a document in its own right. You're doing three things: spotting titles that sound central (harvest them), noticing which references get cited *in the text* at load-bearing moments (the introduction's gap paragraph and the methods' justifications lean on the papers that matter), and watching for sources cited by *everything* you've collected so far. That third signal deserves its own sentence: **when four of your five best papers all cite the same 2011 study, that study is a foundation of your conversation, and you must read it**, whatever its age. Foundations are exactly what date-filtered keyword searches structurally miss.

Forward snowballing is asking who cited a paper since. Scholar's *Cited by* link (Scopus and Web of Science do it with better filters) lists every indexed paper that cites yours. Its uses are sharper than the backward direction:

- **Updating a foundation.** You've found the classic 2011 study; who extended, qualified, or overturned it since? The "cited by" list *is* that history. This is also the what-happened-next check from Chapter 3, mechanized.
- **Escaping a dead end.** An old, slightly-off paper that's not quite your topic often has a recent citer that's *exactly* your topic, because the citer stood where you stand and moved the conversation your way.
- **Searching within citers.** Scholar lets you search *inside* a cited-by list (tick "Search within citing articles"): "papers that cite the foundational study AND mention adolescents" is a devastatingly precise query that no keyword search could express.

One caution carried over from Chapter 3: citation is attention, not endorsement. A heavily cited paper might be heavily criticized (methodological cautionary tales rack up hundreds of citations). When a citation count impresses you, sample a few citing sentences to learn *how* it's being cited: as foundation, as target, or as furniture.

6.2 Reviews and theses as prebuilt maps

Chapter 1 introduced the review article as the conversation summarizing itself; the snowballing frame shows why it's so valuable: **a good review is several hundred hours of someone else's snowballing, organized and annotated.** The optimal use: read the review's structure first (its section headings are a draft taxonomy of your field, and Chapter 17 will steal from them shamelessly), then harvest selectively from its references, then forward-snowball the review itself, because everything that cites a good review on your topic is a candidate for your project.

Recent theses (Chapter 1) work the same way at smaller scale, with a bonus: their literature-review chapters are written by someone at *your* level, so the explanatory walkway is gentler, and their bibliographies were compiled under committee scrutiny.

6.3 Following people and venues

The third strand of the web is social. Conversations have recurring voices, and once your harvested pile has three or four papers, patterns of names appear. When a name shows up on multiple papers you like: visit their Scholar profile (their full output, sorted by year or citations), skim for titles you've missed, and note who their frequent coauthors are; research happens in groups, and finding "the" group for your subtopic can shortcut weeks. The same move works for venues: when three good papers share a journal or conference, browse that venue's recent issues directly. Some corners of every field publish disproportionately in one or two places, and issue-browsing there beats any query.

A light warning on both moves: they concentrate your reading inside one cluster of the conversation, and clusters have house views. If everything you've collected shares authors, venues, or a citation ancestor, you've mapped a neighborhood, not the city, and it's time for one deliberate cross-check search (different database, different vocabulary, or the dissent-hunting move from Chapter 3).

6.4 Citation-graph tools

A newer generation of tools draws the web instead of listing it: you seed them with one or more papers, and they render the citation neighborhood as an interactive graph, clustering related work visually (Connected Papers, Research Rabbit, Litmaps, and kin; Appendix E keeps the current inventory, because this category churns). They're genuinely useful for two jobs: the *orientation moment*, early in a project, when you have three papers and want to see the shape of the room around them; and the *blind-spot audit*, late, when the picture should look familiar, and any large unfamiliar cluster is either irrelevant (fine, note why) or a gap in your coverage (better to find it now than in the viva).

Treat the graphs as suggestions, not authority: their similarity measures are opaque, their coverage uneven. The disciplined loop, tool-agnostic, is: seed, explore, *harvest into the reference manager with a note about why* (Chapter 11 will insist on the why), never browse-and-forget. Twenty minutes of graph exploration that leaves nothing in your database was entertainment, not research.

Keywords find what shares your words. The web finds what shares your ancestors: the foundations everything cites, the frontier that cites everything, and the neighbors you didn't know you had. Every strong paper you find is not a result; it's a new starting node.

6.5 The web session, worked

Continuing Chapter 5's worked project (Instagram and teen anxiety), here's the citation-web session that follows the search session, forty-five minutes again.

Minutes 0–15: backward. Open the harvested review's bibliography beside the big recent study's. Ten minutes of title-scanning yields six harvests, including, three times cited by both, a 2017 longitudinal study your keywords never surfaced (it says “digital screen engagement,” not “social media”). Note the pattern: both bibliographies also cite a 1998 paper on survey methodology; park it, methods foundations can wait.

Minutes 15–30: forward. Cited-by on the 2017 study: 400 citers. Search within citing articles for anxiety: 60. Sort by date, skim two pages: three harvests, one of which is the dissenting camp's flagship (found organically this time). Its own cited-by, quickly checked, reveals the current back-and-forth: two responses and a rejoinder. The live controversy, mapped, in nine minutes.

Minutes 30–40: people and venues. One author appears on four harvested papers; her profile adds two more, including an in-press piece. Two journals recur; a five-minute browse of one's current issue adds a final harvest.

Minutes 40–45: log and audit. Update the search log (sources checked, harvests, the new vocabulary “digital screen engagement”). Quick cluster check: the pile now spans three groups, two methodologies, both camps of the controversy, and one methods classic. That's a real map forming, and tomorrow's reading queue (Chapter 12) just filled itself.

6.6 Managing the snowball

Snowballing has a failure mode of its own: it compounds. Every mined bibliography adds three papers, each of which has a bibliography, and by session's end the harvest queue has doubled twice. Three disciplines keep it an instrument instead of an avalanche:

Harvest with a budget. Before opening a reference list, decide the session's harvest cap (eight to twelve is plenty). Scarcity forces the judgment that matters: not “could this be relevant?” (almost anything could) but “is this more promising than what I'm already carrying?”

Tag provenance. The why-note at capture (Chapter 7) should say who sent you: “via Rev2020 refs,” “cites verbeek2023.” Provenance earns its keep twice: during triage, it's context that speeds pass 1; during writing, it reconstructs intellectual lineage (“the measurement critique originates with...”), which is exactly the kind of sentence that makes a review read as scholarship.

Snowball in sessions, not moods. The citation web rewards focused walks (Chapter 5's session structure applies) and punishes ambient clicking. Twenty minutes of directed snowballing from your best paper is a technique; forty browser tabs at midnight is weather. The search log (Chapter 8) records the former; the latter records nothing and teaches nothing.

Put it to work

- Take your single best paper and mine its bibliography for fifteen minutes. Harvest at least three sources, and note any reference that appears in multiple papers you've already collected; that's a foundation, and it goes to the top of the reading queue.
- Forward-snowball the same paper: open its cited-by list, search within it for your narrowest concept, and harvest what surfaces.

- Identify one recurring author in your pile and spend ten minutes on their profile; one recurring venue, and browse its latest issue.
- Seed one citation-graph tool with your three best papers. Harvest anything genuinely new, with a why-note, and write one sentence about any cluster you're deliberately ignoring.

Getting the Paper

You've found the perfect paper and clicked through to a page demanding \$39.95 for 24-hour access. Do not pay it. Almost nobody who does research for a living has ever paid that number: there is a whole layered system for reaching papers legitimately, and students routinely know only its worst layer. This short chapter is the access toolkit: the legal routes, in the order to try them, plus the thirty-second filing habit that turns each captured PDF from future clutter into future capability.

7.1 Why the wall is there

One paragraph of context, because it explains the workarounds' legitimacy. Traditional journals charge subscriptions; universities pay them (collectively, billions) so their members read freely. The open-access movement, over two decades, has pushed much of the literature out from behind the wall through two main channels: **gold** open access (the published version is free for everyone, the journal financed otherwise, often by author fees) and **green** open access (the journal permits authors to post a version, usually the accepted manuscript, in repositories). The practical consequence: *most papers you need have a legal free copy somewhere or are covered by your library*; the skill is knowing the lookup order.

7.2 The access ladder

Try these in order; the first three resolve the overwhelming majority of cases.

Rung 1: your library's link. On campus networks (or via the library's proxy/VPN off campus), paywalls often simply vanish, because your university subscribes. Install your library's browser extension or bookmarklet (nearly every university has one) so that the "get access via your institution" path is one click from any publisher page. Scholar has a setting (*Library links*) that adds your university's access button directly into search results; turn it on today.

Rung 2: the open copy. If the library lacks the journal, look for the green/gold copy: the **Unpaywall** extension (and its kin) checks legal open repositories automatically and lights up when a free version exists; Scholar's right-hand PDF links do a rougher version of the same. Check the field repository directly (arXiv, bioRxiv, SSRN, PubMed Central) and the authors' university repository or personal pages, where accepted manuscripts legally live. Caveat from Chapter 3: repository copies are sometimes earlier versions; note which version you have, and for anything load-bearing, verify against the version of record before you quote page numbers.

Rung 3: interlibrary loan. For the genuinely unsubscribed paper or book chapter, ILL is the library's superpower: they request it from another library and email you a PDF,

typically within days, free. Students treat ILL as exotic; it's routine plumbing, and the request form takes two minutes. The only cost is latency, which Chapter 12's reading-queue discipline absorbs painlessly: request now, read when it lands.

Rung 4: ask a human. Email the corresponding author (their address is on the paper): two polite sentences, who you are, which paper you'd like. Authors are usually delighted; sharing their own work is legal nearly everywhere and career-positive besides. Response rates are imperfect, latency is days-to-never, but for obscure or old material it works surprisingly often, and it occasionally starts a genuinely useful conversation. The modern variant, asking on academic social platforms, works too. What this rung is *not* for: papers rungs 1–3 already cover; don't spend a stranger's time to save yourself a proxy login.

You'll notice a famous shadow-library site is absent from this ladder. You know it exists; so does everyone. It's unlawful in most jurisdictions, some universities sanction its use, and, decisive for our purposes, **the legal ladder above covers essentially everything a student project needs**. A workflow that depends on infringing sources is a workflow with a structural weakness; build the durable one instead.

7.3 Books, chapters, and theses

Three genre-specific notes. **Scholarly books:** check the library catalog (print and ebook), then ILL for individual chapters, which usually arrive as PDFs; many university-press books also have OA editions now, findable via the DOAB directory. **Chapters in edited volumes** behave like articles for ILL purposes; request the chapter, not the book. **Theses:** most universities' repositories are open; ProQuest's dissertation database (library-subscribed) covers the rest, and rung 4 works charmingly well, because nobody has ever been sad to be asked for their thesis.

7.4 The filing moment

Now the habit that separates a research database from a downloads folder. **The moment a PDF reaches you is the moment it gets filed**, and filing means one thing: it goes into the reference manager (Chapter 13), attached to its metadata record, which the manager renames and stores by its own rules. Never "save to Desktop for now." The thirty-second version of the ritual:

1. Capture the metadata first (the manager's browser button, from the publisher or Scholar page), so the record exists.
2. Attach the PDF to that record (often automatic; otherwise drag it on).
3. Glance at the metadata: right title, right year, right version (preprint or record?). Fix now; wrong metadata caught today costs ten seconds, caught during bibliography-building costs an evening (Chapter 15).
4. If the paper came from snowballing, add the one-line why-note (Chapter 6's rule): "cited by both reviews re: measurement." Future you, staring at 90 PDFs, will bless this line.

That's the entire chapter's discipline: legal ladder down, metadata-first filing up. It compounds quietly: by Chapter 14, when we build the synthesis matrix, every source will already be present, versioned, findable, and annotated with why it entered your world, and none of Part V's writing will ever stop to hunt for a lost PDF.

Access is a solved problem if you use the system: library first, open copy second, ILL third, author fourth. And a paper isn't "gotten" when the PDF downloads; it's gotten when it's filed, attached, verified, and annotated with why it exists in your library.

7.5 The awkward cases

Four sources that resist the ladder, and their workarounds:

The embargoed thesis. Repositories sometimes delay release for a year or two (patents, pending papers). The request button on the repository page works more often than you'd think, and rung 4 (email the author) works even better; new doctors love being asked.

The conference talk with no paper. Some fields post video and slides; cite what exists (the talk, dated, with its venue) and email the speaker for the manuscript; "in preparation" answers are themselves useful intelligence about where the conversation is heading.

The dataset or codebase. Increasingly first-class sources (Chapter 15 covers citing them). Access runs through repositories (Zenodo, OSF, field archives), and "data available on request" has a well-documented non-response problem: budget for it, ask early, and note the outcome either way; a polite unanswered request is worth one honest sentence in a methods discussion.

The old and scanned. Pre-digital material arrives as ILL scans or archive.org copies; verify page numbers against the edition you cite, because scans of different printings disagree, and Chapter 15's page-number discipline needs a stable target.

Put it to work

- Set up rungs 1 and 2 today: library proxy or VPN, the library's browser extension, Scholar's Library links setting, and an open-access checker extension. Fifteen minutes, permanent payoff.
- Find one paywalled paper from your pile and walk the ladder until you have it legally. Note which rung worked; it's usually lower than students expect.
- File an ILL request for the one source you've been postponing (that book chapter counts). The latency clock starts only when you click.
- Audit your downloads folder for stray research PDFs and run each through the filing ritual, or delete it. End at zero strays, and stay there.

8

Knowing When to Stop

Searching has a failure mode nobody warns you about, and it isn't missing papers. It's not stopping. The search is comfortable: it feels like work, produces visible progress (look at all these PDFs!), and postpones the harder acts of reading and writing. Students ride that comfort for weeks. Meanwhile the anxiety runs underneath: *somewhere out there is the paper that either proves my idea is taken or would transform my project, and if I stop looking I'll be exposed in the viva*. This chapter replaces both the comfort and the anxiety with something better: stopping criteria you can actually check, and a written log that converts "I think I looked everywhere" into evidence.

8.1 Saturation: the signal you can feel

The core concept comes from the systematic-review world, where stopping is a methodological requirement: **saturation**, the point where additional searching stops yielding new information. You'll recognize it by three concrete signs:

- **The circle closes.** New searches return papers you've already harvested. When your fourth differently-worded query surfaces the same fifteen core papers, the room has walls and you've found them.
- **Bibliographies stop surprising you.** Backward snowballing (Chapter 6) on a newly found paper yields nothing unfamiliar: you recognize every load-bearing reference. This is the strongest single signal, because reference lists are the conversation's own map of itself.
- **New finds stop changing your picture.** The papers still trickling in are variations: another sample, another country, a minor method twist, slotting into themes you already have (the synthesis matrix of Chapter 14 makes "slotting in" literal and visible). Information is arriving; *news* isn't.

Saturation is always *local*: you saturate a specific question at a specific scope, not "the literature." Which is why stopping is really a scoping decision, and scope is where we go next.

8.2 Scope is a sentence you write down

Most unstoppable searches are unscoped searches. The remedy is embarrassingly simple: **write the boundaries down as sentences, and let the sentences do the refusing**. A scope statement for our running example (Chapters 5–6) might read:

I am reviewing quantitative evidence on the relationship between social-media use and anxiety in adolescents (roughly 12–18), published since 2010 in English, in peer-

reviewed venues plus major preprints. Out of scope: clinical treatment studies, adult populations, general “screen time” without a social component, and the qualitative interview literature (except one review for context).

Five lines, and suddenly a dozen tempting rabbit holes have a posted sign: *not this project*. The scope statement isn’t a cage; you’re allowed to amend it when the literature teaches you something (you’ll amend it consciously, in writing, which is entirely different from drifting). It also becomes free raw material for your review’s introduction and your methods section, both of which need exactly these sentences (Chapter 17).

Scope has a second dimension: **depth per source**. Not everything in scope deserves pass-3 reading (Chapter 9); your scope statement can say, and should, something like “core empirical papers read fully; replications and adjacent work logged at abstract level.” Now “covering the literature” has a defined, finite meaning.

8.3 The search log

The tool that makes all of this auditable is the humble **search log**: a running table, one row per search session, kept anywhere durable (a note in your reference manager works; Appendix A bundles a format):

date	tool	query / action	new finds
06-11	Scholar	("social media" OR ...) ...	7 (2 core)
06-11	Scholar	same, since-2022 filter	4 (1 core)
06-13	PsycINFO	thesaurus systematic pass	3 (1 core)
06-15	snowball	refs of Rev2020 + Big2023	6 (3 core)
06-19	Scholar	rebuilt query w/ new vocab	2 (0 core)
06-22	snowball	refs of new finds	0

Three payoffs, in ascending importance. It **prevents redundancy**: you’ll never rerun a forgotten query or re-mine a bibliography. It **writes your methods section**: for theses (and any systematic-flavored review), “databases searched, terms used, dates” is a required paragraph, and the log *is* that paragraph. And it **shows saturation happening**: look at the new-finds column above collapsing toward zero. That’s not a feeling; that’s data, and it’s the honest answer to the viva’s “how do you know you covered the literature?”, not “I searched a lot” but “here is my search protocol and its convergence.”

8.4 Time-boxing and the returns curve

Layer a resource rule on top of the evidence rules, because projects have deadlines: **searching gets a budget**, and the budget reflects the returns curve. Literature search has steep early returns (the first five sessions find the room, the reviews, the foundations) and a long flat tail. A sane default for a semester-scale project: an intensive front week (three or four real sessions of the Chapter 5 and 6 kind), then **one maintenance session per week or two** for the project’s life, plus the automated alerts of Chapter 12 doing ambient scanning between. A doctoral literature review scales the numbers, not the shape.

The maintenance sessions matter because stopping isn’t permanent: the conversation continues while you work, and Part V (Chapter 18) includes the update loop that folds late-breaking work into a living review. What the budget kills is the daily, anxious, unlogged re-searching that masquerades as diligence while actually deferring the reading queue (Chapter 12) and the writing (Part V).

“Have I found everything?” is unanswerable and therefore endless. “Has my logged, scoped search converged?” is checkable and therefore finishable. Stop when the log says so, schedule the maintenance pass, and go read what you found.

8.5 The permission paragraph

One more thing, because criteria alone don’t dissolve the anxiety. Every thesis that has ever been examined missed papers. Every published review has a scope section that names what it excluded. Every examiner has personally experienced the paper that surfaces the week after submission. **Completeness is not the standard, and no one who matters believes it is.** The standard is: a defensible scope, a systematic search within it, demonstrated convergence, and honest handling of what you found. That standard is entirely achievable, you now have the tools for every clause of it, and the search log is your receipt. Stop gathering. The next part of this book is about what the gathering was for.

8.6 Two case studies in stopping

The over-searcher. A master’s student I’ll call Priya had, by week six, 340 PDFs, no notes, and rising panic; every session found more, and reading felt impossible until the finding was “done.” The intervention was exactly this chapter: a scope statement (which instantly orphaned a third of the pile), a reconstructed log (which showed her last three sessions had found nothing new in scope), and a declared saturation. She read for five weeks, wrote a strong review, and her viva’s literature questions were answered from the matrix, not the pile. The 340 PDFs were never the asset; the convergence evidence was.

The under-searcher. A final-year student, call him Dan, had twelve papers, all from the first Scholar page of one query, and a supervisor’s uneasy sign-off. The examiner’s second question named two foundational studies he’d never seen, both two snowball-hops from his own sources. The repair, done for the resubmission, took one afternoon: backward snowballing his three best papers surfaced both missed foundations plus the review that organized everything; one systematic database pass confirmed convergence. Twelve papers wasn’t the sin; an unlogged, unsnowballed twelve was.

Same lesson from both directions: the log and the scope, not the pile’s size, are what make stopping (and starting to write) defensible.

Put it to work

- Write your scope statement, five to eight sentences including explicit exclusions and a depth rule. Pin it where you search.
- Reconstruct your search log retroactively for the sessions you’ve already done (dates, tools, queries as best you recall), then keep it live from today.
- Check your log’s trajectory: if the new-core-finds column hasn’t started falling, book two more structured sessions; if it’s at or near zero across two sessions and two tools, declare saturation in writing and move to Part III.
- Put the maintenance search in your calendar as a recurring 30-minute slot, and treat searching outside it as the procrastination it usually is.

Part III

Reading

9

The Three-Pass Read

Here's the arithmetic nobody says out loud. Your project has gathered, say, sixty candidate papers. Reading a dense academic paper properly takes three to five hours. Sixty papers times four hours is 240 hours: six full-time weeks, just reading, before a word of your own gets written. The arithmetic is impossible, and every working researcher knows it, which is why *nobody who does this for a living reads every paper the same way*. They read in passes: cheap passes for many papers, expensive passes for few, with explicit decision points between. The canonical version (computer scientist S. Keshav wrote the classic two-page note on it) is the three-pass method, and this chapter makes it your default forever.

9.1 Pass 1: the ten-minute triage

Goal: decide whether this paper matters to your project, and file a structural summary either way. **Time:** five to ten minutes. **Route** (this is Chapter 2's triage route, now with a decision attached): title and abstract, dissected into their moves; introduction's final contribution-paragraph; every section heading; figures and tables with captions; conclusion; a glance down the references for names you know.

Then the decision, written down (Chapter 11 gives the note format): one of **discard** (out of scope after all; say why in one line, keep the line, because you *will* re-encounter this paper and the line saves re-triaging), **park** (relevant but peripheral; abstract-level knowledge suffices per your scope's depth rule, Chapter 8), or **promote** (this needs pass 2). A healthy pile triages roughly into thirds, and a first triage of sixty papers is two or three focused sessions, not six weeks.

Pass 1 has a skill inside it: **honest skimming**. Students feel skimming is cheating and compensate by half-reading everything, which is the worst of both worlds: hours spent, no depth gained, nothing decided. Skimming with a decision at the end isn't cheating; it's triage, the same ethically serious act it is in an emergency room. The papers you discard fund the attention the important ones deserve.

9.2 Pass 2: the comprehension read

Goal: understand what the paper claims and how the evidence works, well enough to explain it to a peer and to place it in your synthesis matrix (Chapter 14). **Time:** about an hour for a typical empirical paper. **Route:** Chapter 2's comprehension route, figures first, then results against figures, then methods at skim-plus, then the framing sections last. You read everything except proofs and the deepest method details, and you *take the structured note* as you go; pass 2 without a note is pass 2 you'll repeat in a month.

Two disciplines define a good pass 2. First, **write the claim in your own words** before copying anyone's phrasing: if you can't state the finding without the authors' sentence, you don't have it yet. Second, **mark what you didn't understand** explicitly (a "?" list in the note). Unresolved marks are not failures; they're the agenda for pass 3, a question for your advisor, or an honest limitation of your reading, and naming them beats the fog of general semi-understanding every time.

Pass 2 ends with another decision: is this paper **core** (it will carry weight in your review; needs pass 3 eventually) or is comprehension enough? For most students most papers stop here, and should.

9.3 Pass 3: the adversarial read

Goal: own the paper: understand it the way a referee does, well enough to defend or attack it in a viva. **Time:** two to five hours, sometimes across days. **Reserved for:** the five to fifteen papers your project genuinely stands on.

The pass-3 stance is different in kind, not just depth: **you attempt to virtually re-do the work**. For an empirical paper: reconstruct the design from the methods, ask what you'd have measured and analyzed, then compare against what they did, and chase the discrepancies. For a theoretical paper: work the proofs or derivations with pencil, at least for the central results. For all papers: check every claim in the abstract against the section that's supposed to support it (the contract audit from Chapter 2), read the supplementary material, and follow the two or three load-bearing references to verify they say what's claimed (Chapter 15 will make that verification an ethical requirement for anything you cite through another paper).

Pass 3 is where Chapter 10's full critical toolkit deploys, and its output is the richest kind of note: strengths, genuine limitations (yours, not just the ones the authors admitted), open questions, and the paper's exact relationship to your project. One honest pass 3 teaches you more about a subfield than ten pass 2s, because you finally see how the sausage is made, including the gap between how research is reported and how it's done.

Reading effort is a budget, and passes are how you spend it deliberately: ten minutes to decide, an hour to understand, an afternoon to own. The failure mode isn't reading too little; it's reading everything at the same mediocre depth and deciding nothing.

9.4 Bailing, rereading, and other permissions

Four permissions the method grants that students don't know they have:

You may bail mid-pass. A pass-2 read that reveals the paper is off-scope reverts to a park note with no guilt; sunk minutes don't obligate future hours.

You may stop understanding at your level. Pass 2 on a paper whose statistics outrun your training means honestly noting "analysis beyond me; conclusions consistent with abstract per my read; flagged" rather than performing comprehension. A student citing such a paper for its headline finding, clearly attributed, is doing it right; pretending to have evaluated the mixed-effects model is doing it wrong.

You may reread. Passes aren't one-shot: a paper often gets pass 1 in week two, pass 2 in week five when its theme heats up, pass 3 in week ten when it turns out to be load-bearing. The note accumulates layers (with dates; Chapter 11). This is the normal life cycle of a core source, not indecision.

You may read slowly at first. Your first five pass-2s in a new field will take two hours each, because you're learning the field's vocabulary and furniture, not just the papers. By the fifteenth, an hour. The speed is a compounding return on the notes and vocabulary you're banking; trust the curve.

9.5 The passes meet the pile

Mechanically, the method runs on the reading queue (Chapter 12 builds it): triage sessions run pass 1 in batches over new harvests; comprehension sessions take promoted papers one at a time; pass 3 gets scheduled like the significant undertaking it is, attached to project milestones (before the proposal, before the methods chapter, before the viva). The queue's state lives in your reference manager's tags (to-triage, parked, core, and friends; Chapter 13), so "what should I read today?" always has a two-second answer.

And a rhythm note from the trenches: **don't batch all triage first, all comprehension second.** Interleave. A week of pure pass 1 produces sixty shallow acquaintances and a queue you dread; mixing two triage batches with three real comprehension reads per week keeps understanding compounding while the pile shrinks, and the pass-2 insights sharpen your pass-1 judgment as you go.

9.6 Pass 1, in real time

Here's an actual pass-1, narrated at speed. The paper: one of the harvests from Chapter 5's session.

0:00–0:45. Title promises a longitudinal design. Abstract, dissected: context move, gap move ("few studies track within-person change"), method (two-wave panel, $n=1,200$), result ("small reciprocal associations"), interpretation. Already suspect it's core for the causality theme.

0:45–2:30. Figures: Figure 1 is the cross-lagged model with coefficients; both directions small, one significant. Table 2 has the descriptives; attrition between waves visible (1,200 → 890). Caption tells me the controls.

2:30–4:00. Intro's final paragraph: three contributions claimed; the third ("first in this region") matters for my empty-column hunt. Headings: standard IMRaD, plus a limitations section, good sign. Conclusion: consistent with abstract, no inflation detected at this altitude.

4:00–5:00. References: recognize six of my core papers; two unknown-but-promising titles harvested (the web grows, Chapter 6). Decision, written: promote – causality theme, panel design, region coverage; watch the attrition. Tag flipped, next paper.

Five minutes, one decision, two harvests, and a flag that pass 2 will chase. That's the whole game, twelve times per session.

Put it to work

- Run pass 1, with the written decision line, on ten papers from your pile today. Time-box: ninety minutes for all ten. Note the rough thirds split you get.
- Pick the clearest promote and give it a real pass 2 tomorrow: figures first, own-words claim, "?" list, structured note. One hour, hard stop, then assess what's left understood.

-
- Nominate, provisionally and in writing, your project's likely pass-3 papers (three to five). Don't read them yet; schedule the first one against your next milestone.
 - Add the queue tags to your reference manager now (`to-triage`, `parked`, `promoted`, `core`) and tag your current pile. The whole system goes live in Chapter 12.

10

Reading Critically

“Read critically,” says every syllabus, and then nobody explains what the verb means. It doesn’t mean finding fault everywhere; a reader who trashes everything is as uncalibrated as one who accepts everything (Chapter 3 said the same about sources). Critical reading means one specific discipline: **separating what was shown from what was said, and judging the distance between them**. This chapter equips that discipline at student level: no statistics degree required, just systematic questions asked in the right order.

10.1 Claims versus evidence: the central move

Every paper contains a ladder of statements, from raw observation to grand implication, and Chapter 2 told you the compartments: results hold observations, discussions hold arguments, abstracts hold advertising. The critical reader’s first move on any claim is locating its rung:

1. **What was actually measured or proven?** (“Students in the spaced condition scored 21% higher at 30 days.”)
2. **What’s the licensed generalization?** (“Spacing improved retention *in this population, on this material, at this delay*.”)
3. **What’s the discussion claiming?** (“Spacing should be adopted across curricula.”)
4. **What’s the abstract selling?** (“Spacing transforms learning outcomes.”)

The gaps between rungs are where your judgment lives. Rung 1 to rung 2 asks about internal validity (was the comparison fair?); rung 2 to rung 3 asks about external validity (does it generalize?); rung 3 to rung 4 asks about spin. A paper can be excellent at rung 1 and reckless at rung 4, and your notes (Chapter 11) should record which rungs you’re buying. When you later cite the paper (Part V), you’ll cite the rung you verified, not the rung that sounds best, and that single habit will put you ahead of a depressing fraction of published citers (Chapter 15 returns to this).

10.2 Interrogating methods without a methods degree

You don’t need advanced training to ask the structural questions, because most methodological weaknesses are visible at the level of design logic:

Compared to what? The foundational question. A study showing a method “works” means nothing without the comparison: against doing nothing? Against the current standard? Against a strawman? A surprising number of claims dissolve right here.

Who was studied, and who’s missing? Sample size (is $n = 12$ carrying a universal claim?), but more importantly sample *composition*: undergraduates standing in for human-

ity, one hospital for healthcare, English-language tweets for public opinion. The licensed generalization (rung 2) is bounded by the sample, always.

Could the design see the effect it claims? Correlational designs can't establish causation, however suggestive the language (watch verbs: "linked to," "associated with" describe correlation; "improves," "reduces" claim causation, and require randomized or otherwise causal designs). Cross-sectional snapshots can't establish direction: does social-media use raise anxiety, or do anxious teens use more social media? Our running example (Chapters 5–6) has exactly this structure at its heart, and the field's live controversy is substantially a fight about design.

Who chose what to report? The quiet killers: outcome-switching (measuring ten things, reporting the two that worked), flexible analysis (trying specifications until one crosses the significance line), and publication bias (the literature over-representing positive results because null results don't get submitted). You can't detect these in a single paper, which is exactly the point: they're why Chapter 3 told you to weight conversations over soloists, and why preregistered studies and replications earn extra trust in your matrix.

10.3 A working reader's statistics

The five concepts that repay their learning cost immediately, stated practically:

- **Effect size beats significance.** "Significant" means "probably not zero," not "big." With enough participants, trivial effects reach significance. Always find the magnitude (the 21%, the 0.2 standard deviations, the 3 points on a 100-point scale) and ask whether it would matter in the world.
- **Intervals beat points.** A 95% confidence interval (the 14–28% in our specimen abstract) shows the precision; an interval spanning "tiny to huge" means the study measured imprecisely, whatever its headline.
- **$p = 0.049$ and $p = 0.001$ are different news,** and a literature whose findings cluster just under the 0.05 line is showing you its selection process, not its truth.
- **Subgroups found after the fact are stories, not findings.** "The effect held for left-handed participants over 40" discovered post hoc is exploration, publishable as a hypothesis, citable as nothing more.
- **Replication is the actual statistics.** One study is one draw. Your confidence should track the pattern across studies, which is what meta-analyses and systematic reviews aggregate, and what your own synthesis matrix (Chapter 14) does at project scale.

That's the whole kit. When a paper's analysis genuinely outruns it, use Chapter 9's permission: flag it, cite carefully at the rung you can verify, and move on.

10.4 Reading against the grain, fairly

Critical reading has a posture problem to manage in both directions. The toolkit above can curdle into cynicism: every study has limitations, so a determined reader can dismiss anything, which conveniently excuses ignoring evidence you dislike (Chapter 3's confirmation-shopping, inverted). The fairness disciplines:

Grade against the possible, not the ideal. The question isn't "is this study perfect?" but "is this about as good as this question allows, and what does it add anyway?" Some questions can't be randomized (you can't assign childhoods) and their literatures are built from converging imperfect designs; dismissing each piece for the flaw the ethics board required is not rigor, it's misunderstanding the genre.

Steelman before you strike. Write the authors' best case in your note before your objections: it keeps your criticism aimed at what they claimed rather than a convenient caricature, and about a third of the time the steelman dissolves the objection.

Use their limitations section as a floor, not a ceiling. Authors disclose the limitations they can afford to; your job is noticing the ones they didn't. But when a limitation is disclosed, credit it: a paper that says "this is correlational and preliminary" being cited as preliminary correlation is the system working.

Let the paper change you sometimes. The point of critical reading is calibrated belief, and calibration moves both ways. If a well-designed study contradicts your prior, the disciplined response is updating, in writing, in your notes: "stronger than I expected; my assumption X now looks shaky." A reader whose views never move isn't critical; they're closed, expensively.

Critical reading is measuring the distance between shown and said, in both directions, fairly. Its output isn't a verdict, "good paper" or "bad paper," but a calibrated entry: what this paper establishes, at which rung, with what caveats, and what it does to your picture of the conversation.

10.5 The critical read in your notes

Everything this chapter produces flows into the note template (Chapter 11 formalizes it; a preview of the critical fields): the claim *at its licensed rung*, in your words; the design in one line ("RCT, n=240, undergrads, 30-day delay"); the two or three limitations that actually matter (not a ritual list of every conceivable flaw); the steelman sentence; and the calibration line: what this paper does to your working picture. Those five fields, filled honestly across your core papers, are the entire intellectual input your literature review needs; Part V is largely the craft of arranging them.

10.6 Evidence in other dialects

The toolkit above speaks quantitative-empirical. Three other evidence dialects you'll meet, and what rigor looks like in each:

Qualitative studies (interviews, ethnography, case studies) trade breadth for depth, and judging them by survey standards is a category error. Their rigor markers: a sampling *logic* (who was chosen and why, saturation claimed and justified), transparency about coding and analysis (who coded, how disagreements resolved), reflexivity (the researcher's position acknowledged), and claims scaled to design (mechanisms and meanings, not prevalence). A qualitative paper claiming "most teenagers feel..." from nine interviews has overreached; one proposing *how* passive scrolling might amplify anxiety, with vivid mechanism candidates, is doing exactly its job, and your matrix wants that mechanism column filled.

Theory papers argue without new data. Evaluate internal coherence (do the conclusions follow?), consistency with established results (does it contradict solid findings, and does it know it does?), and fertility (does it generate testable predictions others actually test?). A theory's citation afterlife (Chapter 6) tells you whether the field found it usable.

Humanities argument runs on sources and readings: the check is whether the evidence is fairly selected (do the quoted passages represent the text, or cherry-pick it?), whether counter-readings are engaged, and whether the archive consulted matches the claim's scope. The steelman discipline transfers whole; the statistics toolkit stays home.

Put it to work

- Take your best pass-2 paper and write its claim ladder explicitly: measured, licensed, discussed, advertised, one line each. Note the widest gap.
- Run the four method questions (compared to what; who's missing; could the design see it; who chose what to report) on the same paper, in writing, ten minutes.
- Find the effect size and its interval for the paper's headline finding. If you can't find either, that's a finding about the paper; note it.
- Pick the paper in your pile you most *dislike* (everyone has one) and write its steelman: three sentences of its best case. Keep them in its note.

11

Notes That Survive

Eight months from now, at 11 p.m., you'll be drafting your literature review and you'll need to know: which of the sixty papers measured anxiety with the GAD-7? Which one had the dissenting result about passive use, and what exactly was its sample? You will not remember. **The entire value of your reading collapses to whatever your notes can answer when you no longer remember the papers.** Highlights can't answer questions. A wall of yellow in a PDF is a note to a person who no longer exists: the you who still remembered why it was yellow. This chapter builds notes that survive: structured, self-contained, attached to the source, and written for a stranger, because eight-months-from-now you *is* a stranger.

11.1 Why highlighting fails, and what works

The research on learning is unambiguous, and it matches every grad student's experience: highlighting feels like engagement and encodes almost nothing. It marks *recognition* ("this seems important") without requiring *retrieval* or *reformulation*, and it stores nothing searchable, nothing comparable, nothing usable at synthesis time.

What works is generation: **writing claims in your own words, connected to your own project.** The act of reformulating is where comprehension gets tested (Chapter 9's own-words discipline) and the written product is what synthesis (Chapter 14) and drafting (Part V) actually consume. Marking up the PDF while reading is fine as scaffolding, an intermediate step; the note is the deliverable, and a reading session that ends at highlights is a session whose value is set to evaporate.

11.2 The paper note template

One note per source, one format for all of them, kept in your reference manager attached to the item (Chapter 13 covers where; here's the what). Appendix A reproduces this with a fully worked example; the fields:

```
CITEKEY: verbeek2023passive          STATUS: pass 2 (2026-06-18)
IN ONE LINE: what the paper is, for a stranger
CLAIM(S): the licensed findings, my words, with magnitudes
DESIGN: one line (who/what, n, method, comparison)
EVIDENCE NOTES: key results, figures/tables worth returning to
LIMITS THAT MATTER: 2-3, mine as well as theirs
STEELMAN: their best case in one sentence
QUOTES: verbatim + page numbers, sparse, marked ""
MY REACTION: what this does to my picture; disagreements; "?" list
CONNECTS TO: [citekeys] agrees/contradicts/extends/uses-method
FOR MY PROJECT: where this will matter (theme, section, role)
```

Field-by-field craft, the parts students get wrong:

Claims, in your words, with magnitudes. Not “spacing helps memory” but “spaced practice improved 30-day retention by about a fifth (21%, CI 14–28) in undergrads, no sleep interaction.” The rung discipline of Chapter 10 lives here: record what was *licensed*, not what was advertised.

Quotes are rare and ceremonial. Copy verbatim only sentences you might actually quote (a definition, a strikingly put position), always in quotation marks, always with a page number. This is your plagiarism firewall: eight months from now, the marks are the only thing telling you which words are yours and which are theirs, and accidental patchwriting, source sentences half-absorbed into “your” prose, is the most common integrity failure in student reviews (Chapter 15 finishes this argument).

My Reaction is where the value concentrates. Your agreement, your doubt, the connection to your project, the question you’d ask the authors: this field is the difference between a summary service and a thinking record. Write it in first person, informally; it’s the seed of the sentences that will make your literature review *yours* (Chapter 18).

Connects To is the synthesis engine. Every time you finish a note, spend ninety seconds asking: what does this agree with, contradict, extend, or share methods with, among papers I’ve already noted? Those links, citekey to citekey with a verb, are the raw material of themes; when you build the synthesis matrix in Chapter 14, this field will have half-built it already.

11.3 Scaling the effort to the pass

Notes scale with reading depth (Chapter 9), and pretending otherwise kills the habit. The working proportions:

- **Pass 1 note** (every triaged paper, two minutes): citekey, status, the one-liner, and the decision with reason: “discard: adults only, out of scope” / “park: replication of verbeek2023, adds Norway sample” / “promote: dissenting camp flagship.” That single discard line will save you from re-triaging the paper when it resurfaces in another search, which it will.
- **Pass 2 note** (promoted papers, fifteen minutes inside the hour): the full template at working depth: claims, design, two limits, reaction, connects-to. This is the standard note, and most of your core corpus stops here.
- **Pass 3 note** (the load-bearers, grows over sessions): everything above, deepened, plus the adversarial findings: verification of load-bearing citations, your reconstruction differences, open questions. Dated layers as rereads accumulate (“2026-08-02, pass 3 for methods chapter:…”).

The proportions matter because the failure mode of note-taking systems is perfectionism: the beautiful two-hour note on paper four, the guilt, the abandonment of the system by paper nine. A scruffy complete field beats an elegant empty one, every time. Write badly, capture everything that matters, move on.

11.4 Notes are for retrieval, so write for search

Mechanics that make the corpus answer questions later:

- **Anchor on the citekey** (verbeek2023passive), the same key your reference manager and your eventual citations use (Chapters 13 and 15). Notes, PDFs, matrix rows,

and draft citations all join on this key; it's the primary key of your whole research database.

- **Use consistent theme words**, and harvest them into a controlled list as themes emerge ("measurement," "passive-active," "causality-debate"). Consistent words make the corpus greppable: *show me everything tagged measurement* is exactly the question Chapter 14 will ask.
- **Make each note self-contained**. The test: a stranger with only this note (not the paper) could place the source correctly in your review. That stranger, again, is you in month eight.
- **Date everything**: notes record what you thought when; your calibrations will change (Chapter 10 told you to let them), and the dated trail is how you'll see your own understanding move, which is worth seeing.

A note is a message to a future stranger who must act on your behalf. Structured, self-contained, own-words, source-anchored: if the stranger can write your literature review from the notes alone, the notes are done. That is, in fact, the plan.

11.5 One worked note

The template, filled, at honest pass-2 scruffiness (the polished version is in Appendix A):

```
CITEKEY: orben2019screens          STATUS: pass 2 (2026-06-19)
IN ONE LINE: huge-dataset reanalysis arguing screen-tech
              associations w/ teen wellbeing are tiny once analysis
              flexibility is controlled
CLAIM: association exists but is very small (comparable to
      "eating potatoes" in their famous comparison); analysis
      choices swing results more than the effect itself
DESIGN: spec. curve analysis over 3 big cohort datasets,
      n ~ 350k, correlational
LIMITS THAT MATTER: correlational (they say so); wellbeing
      measures self-report; "screen time" lumps very different uses
STEELMAN: even taking measurement seriously, no analysis path
      yields the large effects public discourse assumes
QUOTES: "" none yet
MY REACTION: this reframes my whole Q -- maybe the fight is
      about effect SIZE not existence? ? = how do the panel studies
      post-2019 respond to this?
CONNECTS TO: contradicts twenge2018 (magnitude, not sign);
      uses-method spec-curve -> check simonsohn cite
FOR MY PROJECT: core, causality-debate theme, will anchor the
      "how big is it" section
```

Ten minutes of typing, and notice what exists now that didn't before: a retrievable, comparable, project-linked record that already knows which section of your review it belongs to. Sixty of these are not sixty summaries; they're a database, and Part IV is about making the database hum.

11.6 Notes on non-article sources

The template flexes for the other genres in your pile:

Books and monographs: one note per *chapter* you actually use, keyed author-year-word plus a chapter marker, each with its own claims and page-anchored quotes; a thin parent note holds the book's overall argument and which chapters you read (honesty at citation time, Chapter 15).

Talks and videos: timestamped instead of paged ("claims replication at 31:40"), with a link and the date accessed; treat claims as personal-communication grade until the paper version lands.

Datasets: the note records provenance (who collected, when, instrument), version, license, and known quirks; the LIMITS field earns its keep here more than anywhere.

Personal communications (the email from the author, the hallway answer at a conference): dated, quoted precisely, and marked for permission, since citing private communication generally requires asking. The note exists so that eight months later "someone told me" has a name, a date, and a sentence you can actually use.

Put it to work

- Copy the template into your reference manager's note field (or wherever your notes will live) as a snippet you can paste. Appendix A has the long form.
- Write one real pass-2 note today on your best paper, time-boxed to fifteen minutes, scruffiness encouraged. Fill every field with something, even "none yet."
- Retrofit pass-1 notes (the one-liner and decision line only) onto the last ten papers you triaged. Twenty minutes total; resist depth.
- Start your theme-word list with the first three that suggest themselves from your notes so far. Pin it next to the scope statement from Chapter 8.

12

The Reading Queue

Parts II and III have given you engines: searching that harvests papers, passes that consume them at the right depth, notes that bank the value. What’s missing is the drivetrain connecting them: the system that decides *what gets read when*, keeps the harvest from burying you, and keeps you current without keeping you anxious. That system is the reading queue, and it’s the lightest machinery in this book: a few tags, a weekly rhythm, and two automations. This chapter assembles it and hands you the schedule.

12.1 The queue is tags, not a place

Resist the urge to build anything: no spreadsheet of shame, no elaborate kanban. Your reference manager (Chapter 13) already holds every item; the queue is just a handful of status tags on those items, matching the passes of Chapter 9:

to-triage	fresh harvest, no pass 1 yet
parked	triaged; abstract-level is enough (with reason)
promoted	triaged; awaiting pass 2
core	pass-2 done and load-bearing; pass 3 candidates
discarded	triaged out (with the one-line why); kept, not deleted

A saved search or smart folder per tag, and “what should I read today?” becomes a two-second glance: the to-triage folder is your inbox, promoted is your reading list, core is your project’s spine. Every item’s journey is visible, nothing relies on memory, and the whole apparatus cost five minutes to set up (you did most of it in Chapter 9’s exercises).

Two rules keep the tags honest. **Everything harvested gets to-triage at capture** (part of the filing ritual, Chapter 7), so the inbox is complete by construction. And **no item leaves a state without its artifact**: leaving to-triage requires the pass-1 decision line; leaving promoted requires the pass-2 note (Chapter 11). The tags aren’t aspirations; they’re receipts.

12.2 The weekly rhythm

Reading at project scale is a pacing problem: sprints produce burnout and retention loss; pure responsiveness produces a queue that only grows. The sustainable shape, sized here for a semester-scale project alongside coursework, is a **weekly reading block structure**:

- **One triage session** (60–90 minutes): batch pass 1 over the to-triage inbox, newest harvest first. Target: inbox to zero, or near it, weekly. Ten papers an hour is the practiced rate; your first sessions will be slower.

- **Two to three comprehension sessions** (an hour each): one promoted paper per session, full pass-2 protocol, note and all. Deep work; calendar them like classes, because they are.
- **The maintenance search** (30 minutes, from Chapter 8): the queue's supply line, feeding next week's triage.

That's four to five hours weekly, and the arithmetic works: a sixty-paper corpus triages in the first fortnight and comprehends its promoted third over eight to ten weeks, exactly the shape of a semester project, with pass 3s scheduled against milestones on top. Doctoral scale multiplies the block count, not the design.

The rhythm's quiet superpower is **interleaving** (Chapter 9 previewed it): triage teaches you the field's surface, comprehension its depth, and each sharpens the other week by week. Students who batch ("I'll read everything in January") lose both compounding effects and most of January.

12.3 Staying current without drowning

The conversation continues while you work (Chapter 8), and the queue absorbs it through two automations, set up once:

Alerts on queries. Scholar (and Scopus, and most databases) will email you new results for a saved query. Create alerts for your two or three best queries, the mature ones from your concept table (Chapter 5), plus, in preprint-native fields, the repository's category feed or listing service. Aim the messages at a folder, not your inbox's front page.

Alerts on ancestors. Scholar can also alert you to *new citations of a specific paper*. Set this on your three to five foundations (Chapter 6 found them): anyone who cites the foundational study is, by construction, joining your conversation, and this alert catches them regardless of vocabulary. This is forward snowballing, automated, and it's the single highest-precision current-awareness tool you can own.

Then the discipline: **alerts feed the inbox, not your attention**. Once a week, during triage, skim the alert folder, harvest the plausible (filing ritual as always), let the rest go. An alert is not an obligation; it's ambient radar. If a week's radar yields nothing, delightful, that's saturation holding (Chapter 8), and the maintenance session just got shorter.

12.4 Backlog mercy

Every real queue develops a sediment: the promoted papers that three consecutive weeks didn't reach. The sediment's message is almost never "work harder"; it's recalibration data, and the queue's monthly ritual is reading it. Once a month, skim the whole promoted list and re-decide, coldly: anything you keep skipping is telling you it belongs in parked (be honest about why it felt important; the reason often reveals it isn't), anything three months stale gets re-triaged against your current, sharper scope, and anything that survives re-decision gets *next week's first comprehension slot*, jumping the line.

The corresponding permission, stated plainly because students need it in writing: **an unread backlog is not a debt, and the queue is not a conscience**. It's a decision system, and a paper deliberately parked is a decision made, not a failure logged. The failure mode this chapter exists to prevent is the other one: sixty half-urgent PDFs, no states, no rhythm, and reading time allocated by guilt at midnight. That system, the default system, is the one that doesn't survive contact with week eight.

The queue's product isn't finished papers; it's allocated attention. Tags make the state visible, the weekly rhythm makes progress automatic, alerts make currency ambient, and the monthly mercy pass keeps the whole thing honest. Reading stops being a mood and becomes a schedule.

12.5 The whole pipeline, running

Pause and look at what Parts II and III have assembled, because it's now a complete working system. A paper's life: an alert or search session surfaces it (Chapters 5, 6, 8); the access ladder retrieves it and the filing ritual banks it, tagged to-triage (Chapter 7); the weekly triage gives it pass 1 and a fate (Chapter 9); if promoted, a comprehension session gives it pass 2, a structured critical note, and connections (Chapters 10, 11); if core, it awaits its milestone pass 3. Every stage leaves an artifact; nothing depends on memory or mood; and the accumulating output, a corpus of linked, critical, retrievable notes, is precisely the input Part IV organizes and Part V writes from. The fog from the preface is, at this point in the book, structurally impossible.

12.6 Rhythms for different seasons

The weekly block above assumes a steady semester. Real projects have seasons, and the queue bends without breaking:

Coursework-heavy weeks: protect the triage session first (it's the cheapest and keeps the inbox honest), let comprehension drop to one, and log the debt; the monthly mercy pass will re-balance.

Research summers invert the ratio: triage stays weekly, comprehension doubles, and pass-3 sessions get mornings, because deep adversarial reading is the first casualty of a tired brain.

Writing months (Part V in full swing): reading narrows to *targeted* pulls, the specific paper a drafting session needs, plus the alert skim. Resist the comfort-read of new material as drafting avoidance; the freeze discipline of Chapter 18 exists for exactly this season.

The final fortnight: reading stops except for load-bearing arrivals, and the queue's only job is the citation pass of Chapter 15. An unread inbox at submission is not a failure; it's the next project's head start.

Put it to work

- Finish the tag system and saved searches in your reference manager (five minutes; Chapter 9 started it), and drive your current pile into honest states.
- Put the weekly block in your actual calendar: one triage, two comprehension, one maintenance search. Recurring, named, defended.
- Set up both alert types today: two query alerts, plus citation alerts on your three foundations. Route them to a folder.
- Run the first monthly mercy pass right now on whatever backlog you already have: re-decide every stale item, park without guilt, and give one survivor next week's first slot.

Part IV

Organizing

13

Your Reference Manager

Every chapter so far has quietly assumed a piece of infrastructure: somewhere that holds your sources, their PDFs, their notes, their tags, and later, their citations. That somewhere is a reference manager, it's the single most important tool purchase of your student career, and the purchase price is zero. This chapter sets one up properly, end to end, so that every ritual this book has taught (filing, tagging, noting) runs frictionlessly, and so the bibliography of every document you write for the next decade generates itself.

13.1 Zotero, and the one-page honest alternatives

This book's default is **Zotero**: free, open-source, nonprofit, cross-platform, with the best browser capture in the business, built-in PDF annotation, plain-good notes, tags, saved searches, group libraries, and plugins for everything (including the Better BibTeX plugin that LaTeX workflows lean on, per my *LaTeX for Theses and Dissertations*). When this book says “reference manager,” picture Zotero.

The honest alternatives, one line each, expanded in Appendix E: **Mendeley** (capable, but Elsevier-owned with a shrinking feature set; hard to recommend new adoptions), **End-Note** (paid, institutional, heavyweight; use it if your lab standardizes on it), **Paperpile** (slick, subscription, Google-ecosystem native; lovely if you live in Docs), **JabRef** (open-source .bib-native; a fine choice for pure-LaTeX lives). Switching costs are real but survivable (everything exports); the important decision isn't which manager, it's *one manager, used properly, starting now*. Students who delay this setup until writing time do their bibliography by hand once, badly, at midnight, and convert instantly.

13.2 Setup, in one sitting

Forty-five minutes, once:

1. **Install** the desktop app and, crucially, the browser **connector** extension. The connector is the capture ritual's engine: one click on any publisher page, Scholar result, arXiv listing, or news article, and the item lands in the library with metadata and (when accessible) the PDF. This is the “capture the metadata first” step of Chapter 7, mechanized.
2. **Create an online account and turn on sync.** Your library now lives on every machine you use and survives any of them dying. Attachment storage above the free tier: either pay the modest fee (the simplest backup money you'll spend) or wire the WebDAV/linked-file alternatives (Appendix E); either way, **the library must not exist on exactly one laptop**, for reasons Chapter 14 of my thesis book states grimly: one machine will eventually have a bad day.

3. **Set the citation-key format** (via Better BibTeX if you'll ever touch LaTeX; its key pattern setting) to the convention your notes already use: author-year-word, orben2019screens. One key joining manager, notes, matrix, and manuscript is the primary-key principle of Chapter 11, and it gets configured here.
4. **Make the skeleton:** collections (folders) per project, the status tags of Chapter 12 with their saved searches, and a `_inbox` default landing collection so captures never vanish into the void.
5. **Point your word processor at it:** the Word and LibreOffice plugins install with the app; Google Docs works via the connector. Citing becomes a picker dialog (Chapter 15 runs the workflow).

13.3 Metadata hygiene

The manager is a database, and databases repay cleanliness. The capture connector gets metadata right most of the time; your job is the thirty-second verification at filing (Chapter 7) plus a few standing rules:

- **Item type matters:** a conference paper captured as a journal article cites wrong in every style forever. Fix type at capture, especially for the genres of Chapter 1 (preprints! theses! reports!), which capture most raggedly.
- **Titles in one casing convention** (the manager can store sentence case and let styles handle the rest), authors as structured names not text blobs, years present, DOIs present wherever they exist: the DOI is the item's global primary key and powers every dedupe and update operation later.
- **Versions get noted:** the preprint-vs-record distinction of Chapters 3 and 7 lives in the item (Zotero's Extra field, or the version fields where supported), so citation time doesn't rediscover the question.
- **Duplicates get merged on sight:** the built-in duplicate detector, run during the weekly triage session, keeps snowballing's re-finds (Chapter 6) from fragmenting your notes across twin records.

None of this is perfectionism; it's the difference between a bibliography that assembles itself correctly (Chapter 15) and an evening of hand-fixing broken entries during submission week, the exact evening my thesis book's bibliography chapter tries to spare you.

13.4 Organizing: collections, tags, and search

Three organizing tools, three distinct jobs, and most students misuse at least one:

Collections are project views. An item can live in many collections at once (they're playlists, not folders-with-copies); make one per project or chapter, and don't taxonomize the universe: deep collection trees ("Psychology > Social > Media > Adolescents > Anxiety > Longitudinal") are a procrastination format. Two levels, project-shaped, is plenty.

Tags are cross-cutting states and themes. The status tags of Chapter 12 (workflow state) and the theme words of Chapter 11 (content), exactly as those chapters built them. Keep the vocabulary controlled: harvest new theme tags deliberately into your pinned list, prune synonyms monthly, and the tag system stays a query language instead of becoming a compost heap.

Search is the actual interface. Saved searches for the queue states; full-text search across PDFs and notes for everything else ("which paper mentioned GAD-7?" is a three-second query over a well-fed library, the exact 11 p.m. question this part of the book exists to

answer). The manager’s search only sees what you put in it, which is the final argument for the filing ritual: **a source outside the manager is invisible to every future question.**

13.5 Groups, sharing, and the lab

Zotero’s group libraries give a shared collection to a lab, a co-authored project, or a reading group: shared items, shared PDFs (license permitting), shared tags. Two etiquette rules cover most collaborative grief: agree on the metadata and tagging conventions before populating (send them this chapter), and keep personal reading states personal: your parked is not a group verdict. For supervised students, a shared project collection with your advisor is a quietly excellent practice: “have you seen X?” becomes drag-and-drop instead of email attachments, and your growing, well-kept library is itself visible evidence of a project under control.

The reference manager is the single source of truth for everything you’ve gathered: meta-data, PDFs, notes, states, themes. One tool, one key scheme, everything captured at arrival, hygiene at the edges, and every later stage of the workflow, synthesis, citation, writing, inherits a clean database instead of a shoebox.

13.6 Migration and maintenance

If you’re arriving with an existing mess (a folder of PDFs, a half-used Mendeley, browser bookmarks), migrate in one bounded session, not as a background guilt: bulk-import the PDFs (Zotero retrieves metadata for most via their embedded DOIs), import the old manager’s export file, merge duplicates, drive everything through a fast triage into the tag system, and delete the old shrines. Perfection is not required; *closure* is. From then on, maintenance is already scheduled: dupes and inbox during weekly triage (Chapter 12), tag pruning and the mercy pass monthly. The library becomes like a well-run kitchen: cleaned as you cook, never a special project again.

13.7 Beyond papers: what else lives in the library

A research library that holds only journal articles is missing half the conversation (Chapter 1’s genre lesson, applied to storage). Capture, with correct item types: **theses** (type: thesis; the repository page captures cleanly), **books and chapters** (separate item types, and chapter items carry their page ranges into citations), **datasets and software** (modern types exist for both, with version fields your methods section will need), **reports and standards** (the gray literature of Chapter 3, with issuing body as “institution”), **talks** (with URL and date), and **web pages** (with an archived snapshot: pair the capture with an archive.org save, paste the archived URL into the record, and future-proof the citation against link rot, which afflicts a third of web citations within a few years). The rule underneath: *if you might cite it, it gets a record*; the manager is the single source of truth, and truth includes the awkward genres. What stays out: the merely-interesting (bookmarks are not a library), and anything you haven’t at least pass-1 processed, which the inbox discipline of Chapter 12 already enforces.

Put it to work

- Do the five-step setup sitting today: app, connector, sync, key format, skeleton, word-processor plugin. Forty-five minutes, once, forever.
- Run the migration session on your existing pile, bounded to two hours, ending in honest tags and zero orphan PDFs (Chapter 7's audit, finished for good).
- Adopt the three hygiene reflexes at capture: type check, DOI present, version noted. Say them out loud thrice; they're yours now.
- If you work with a lab or advisor, propose the shared collection this week, conventions attached.

14

The Personal Research Database

Sixty sources, sixty structured notes, clean metadata, honest tags: Part III and Chapter 13 built you a collection. This chapter turns the collection into an *instrument*: something you can ask questions of, see patterns in, and eventually write from. The centerpiece is the synthesis matrix, the single most underrated tool in student research, and the bridge on which the entire literature review (Part V) will cross from “what I read” to “what I have to say.”

14.1 From notes to sight

Here’s the problem the matrix solves. Your notes are deep but *serial*: one paper at a time, each thorough, none comparative. A literature review, though, is made of *lateral* claims: “most studies find X,” “the field splits on Y,” “only two papers address Z.” Lateral claims require seeing across papers simultaneously, and no stack of serial notes, however good, provides that sight line. Students without a lateral tool write the book-report review, paper after paper after paper, because serial notes are all they can see. The matrix is the lateral tool.

14.2 Building the synthesis matrix

A synthesis matrix is just a table: **sources down the rows, themes across the columns, and in each cell, what that source says about that theme** (or a dash: silence is data too). Build it in anything rectangular: a spreadsheet is perfect (Appendix B ships a worked one).

Rows are your core corpus: the core-tagged papers (Chapter 12), anchored by citekey, plus a few columns of standing metadata (year, design one-liner from the note’s DESIGN field).

Columns are themes, and choosing them is the first act of synthesis. You don’t invent them; you harvest them. Three sources, in order of yield:

1. **Your Connects-To fields** (Chapter 11): the agree/contradict/extend links you’ve been writing cluster naturally around a handful of fault lines; each cluster is a theme.
2. **Your theme-word list** (already controlled, already pinned): mature tags are column candidates by definition.
3. **The reviews’ section headings** (Chapter 6 told you to study their structure): the field’s own taxonomy, ready to steal and adapt.

For our running project, the columns that emerge look like: *effect size found*; *design (cross-sectional, longitudinal, experimental)*; *measurement of use (self-report vs logged)*;

passive-vs-active distinction; population; stance in the causality debate. Six columns, and notice they're *questions you'd ask every paper*, which is exactly what a theme is.

Cells are one or two lines, compressed from the note's CLAIM field, in your words, with magnitudes where they matter ("small assoc. ($r \approx .05$), spec-curve, logged data" beats "finds effect"). Filling a row takes five minutes per paper because Chapter 11 already did the thinking; you're transcribing, laterally.

14.3 Reading the matrix

Fill twenty rows and the patterns start shouting. This is the payoff, and it's worth naming the four kinds of pattern, because each becomes a different kind of sentence in your review (Chapter 17 completes this handoff):

- **Consensus columns:** cells mostly agreeing. "Nearly all studies find a small positive association" is now a sentence you can write *and defend with a column*.
- **Fault-line columns:** cells splitting into camps, and, run your eye to the design column, often splitting *with* a methodological correlate: the big effects cluster in cross-sectional self-report studies; the tiny ones in logged-data designs. That correlation between finding and method is the single most valuable kind of insight a literature review can offer, and the matrix makes it visible mechanically.
- **Empty columns:** a theme almost nobody addresses. "Only two studies examine logged rather than self-reported use": that sentence, defensible from the matrix, *is* a research gap, and Chapter 17 will teach you to state it honestly (a gap the field ignored on purpose is different from one it hasn't reached).
- **Empty rows:** a core paper silent on most themes. Either it's less core than you thought (demote it, Chapter 12's mercy applies) or your columns miss what it's about (add the column it deserves; the matrix is allowed to grow).

There's a fifth pattern, personal: the matrix shows you *your own coverage*. A column you can't fill for half your rows means your notes lack that field, which means your reading passes weren't asking that question, which is fixable next session. The instrument calibrates its operator.

14.4 Keeping it living, keeping it light

Process notes, learned from many student matrices:

Start at fifteen rows, not sixty. The matrix is for your core corpus; feeding it every parked paper turns an instrument into a data-entry job. Parked papers contribute through their one-line notes when a specific cell needs them.

Update on pass, not on schedule. A new pass-2 note ends with its matrix row (five minutes); a new theme emerging mid-project gets its column retrofitted in one dedicated half-hour, not paper-by-paper as guilt strikes.

Version it and date it like everything else; the matrix at proposal time versus viva time is a record of your understanding maturing, and more practically, a stale matrix misleads confidently.

Two tools, one join key. The matrix never contains the full notes, just compressions plus the citekey, which is the join back to everything deeper (Chapter 11's primary-key principle at work). Spreadsheet for the lateral sight, manager for the vertical depth, citekey between them: that's the entire architecture of the personal research database, and it's two free tools and a naming convention.

Notes let you know each paper; the matrix lets you see the field. Consensus, fault lines, and silence become visible, checkable, and citable, and the literature review stops being an essay you must conjure and becomes a description of a table you already possess.

14.5 Beyond the matrix: when you need more

For most student projects the matrix plus the manager is the whole database, and adding machinery subtracts momentum. Two genuine step-ups exist, flagged honestly (Appendix E details tooling):

Linked note systems (Obsidian and kin, the “Zettelkasten” tradition): notes as a network with explicit links, superb for multi-year, multi-project intellectual life, where themes recur across papers, projects, and years. The principles transfer directly (own words, one idea per note, links as first-class content, the citekey as anchor); the overhead is real, and a thesis deadline is the wrong week to adopt a note philosophy. If your Connects-To fields are bursting and your projects span years, explore it between projects.

Extraction tables for systematic reviews: if your field’s norms (medicine, increasingly others) require a formal systematic review, the matrix grows into a preregistered extraction protocol with defined fields, dual coding, and PRISMA-style flow accounting. That’s this chapter’s machinery with regulatory plating; the concepts are identical, and your matrix habit is the on-ramp.

14.6 Matrix pathologies

Four ways matrices go wrong, each visible and fixable:

The double column: cells keep wanting two answers (“effect size” answers split by short/long term). The column is two columns; split it, and the fault line you were blurring becomes visible, which is usually the insight.

The misfit row: a theory paper or review that resists your empirical columns. Don’t force claims into cells built for findings; give the genre its own columns (“framework proposed,” “predictions offered”) or let it live in the notes and feed the review’s framing instead. Not everything lateral is tabular.

Cell bloat: paragraphs creeping into cells until the lateral sight line dies. Cells compress; notes hold depth; the citekey is the door between them. If a cell can’t compress, the paper probably deserves a pass-3 and the cell a humbler summary.

The private matrix: kept, but never shown. The matrix is your best advisor-meeting artifact: five minutes of walking the columns beats twenty of “how’s the reading going,” and an advisor’s “you’re missing the X literature” lands as a concrete new column instead of a vague worry. Bring it; version it; let it be seen.

Put it to work

- Create the matrix today: spreadsheet, citekey rows for your current core papers, year and design columns, then harvest four to six theme columns from your Connects-To fields, tag list, and the best review’s headings.
- Fill five rows from your existing notes, timed: it should run about five minutes each. If a cell requires reopening the PDF, your note missed something; improve the note, then the cell.

- Read the half-filled matrix for ten minutes and write down, in sentences, one consensus, one fault line, and one empty column you can already see. Date the page; it's the first draft of your review's spine.
- Add "matrix row" as the closing step of your pass-2 ritual, right after the note's FOR MY PROJECT field.

15

Citing Without Tears

Citation is where your research database finally touches the documents you submit, and it's a strange double subject: half mechanical (styles, formats, tooling) and half ethical (what citing claims, and what miscounting it costs). Students fear the mechanical half and underestimate the ethical half; the truth is inverted. The mechanics, after Chapter 13, are nearly free. The ethics are where marks, and occasionally degrees, are actually lost. This chapter settles both.

15.1 What a citation claims

Start with meaning, because everything else follows from it. A citation attached to a claim asserts three things:

1. *This claim traces to this source* (attribution),
2. *the source actually supports it, at this strength* (verification), and
3. *I have engaged with the source* (honesty about your process).

Every citation sin is a failure of one of the three. Misattribution fails the first. Citing a paper for more than it licensed, the rung inflation of Chapter 10, fails the second, and it's endemic: follow citations in published work (you did, in Chapter 2's claim-checking route) and you'll find a startling fraction misrepresent their sources, usually by inheriting someone else's citation without reading. And that's the third failure: **citation laundering**, citing what you never read because someone else cited it. The professional norm, and your standing rule: **cite only what you have at least pass-1 read, at only the strength your note supports**. When you genuinely need a source you can't access (the untranslated 1962 classic), the honest form is the secondary citation: "(Dubois, 1962, as cited in Chen, 2021)," which tells the reader exactly whose reading you're trusting. Use it sparingly; each one is a small confession that the ladder of Chapter 7 came up short.

15.2 Quotation, paraphrase, and the patchwriting trap

Three ways to bring a source's content into your prose, with sharply different rules:

Quotation: their words, quotation marks, citation with page number. In most of the sciences, quotation is rare and should be: you quote definitions, coinages, and positions where the wording itself matters. In the humanities, quotation is the evidence itself and flows constantly. Know your field's dialect (Chapter 2's advice, again).

Paraphrase: their idea, your words, citation still required, always. Paraphrase is the workhorse of scientific prose, and its standard is *genuine* rewording: your sentence structure, your vocabulary, the idea passed through your understanding (which is why Chap-

ter 11 made own-words claims the note's core field: your notes are pre-laundered legitimate paraphrase, written months before deadline pressure).

Patchwriting is the trap between them: the source's sentence with synonyms swapped and clauses shuffled, no quotation marks. It's how most student plagiarism actually happens, not intent, but drafting-from-the-PDF at midnight, half-absorbing sentences that were never yours. Integrity software flags it, committees treat it as plagiarism, and the defense "I changed the words" has never once worked. The structural fix is the one this book already built: **draft from your notes, never from the source**, with quotes ceremonially marked in the notes themselves (Chapter 11's firewall). Writers who draft from notes cannot patchwrite, because the source's sentences aren't on the desk.

Self-plagiarism rounds out the ethics: reusing your own graded or published text without disclosure is an integrity violation at most institutions (yes, really; the reader is being told this is new work). Reusing your own *ideas* is called a research program and is encouraged. When in doubt, disclose and ask; the question costs nothing.

15.3 Styles, and what they're actually doing

Citation styles (APA, MLA, Chicago, IEEE, Vancouver, and the thousands of journal variants) are surface conventions over identical information. The families:

- **Author-date** (APA, Chicago author-date, Harvard): "(Orben & Przybylski, 2019)" inline. Dominant in social and life sciences; keeps authors and recency visible in the prose, which suits conversations where *who* and *when* carry meaning.
- **Numeric** (IEEE, Vancouver, most physical sciences): "[12]" inline. Compact, suits dense multi-citation in constrained pages.
- **Notes** (Chicago notes-bibliography): footnoted citations; the humanities' home, where sources need commentary.

Two practical consequences. First, **you don't choose; you comply**: the journal, department, or field chooses, and the choice is in the syllabus or author guidelines. Second, **you never hand-format**: your manager holds the information; the style is a rendering. Chapter 13 installed the word-processor plugin: you cite via picker, the bibliography assembles and reorders itself, and switching a document from APA to Vancouver is a dropdown, not a weekend. (LaTeX users: the same separation is biblatex's whole design, and my thesis book's bibliography chapter runs that side.) Your only formatting jobs are the ones automation can't do: choosing *where* citations attach (next section) and the final proofread, since style rendering is only as clean as the metadata hygiene of Chapter 13, which is why that hygiene exists.

15.4 Citation craft in the sentence

Where and how citations sit in prose is a craft with real information content, and it previews Part V's writing:

Attach citations to claims, not paragraphs. A citation at a paragraph's end vaguely blesses everything above it; a citation after each specific claim tells the reader exactly what rests on what. When one sentence carries three sourced claims, it carries three citations, at their clauses.

Choose narrative or parenthetical deliberately. "Orben and Przybylski (2019) found..." foregrounds the actors: use it when the *who* matters (camps, debates, methods traditions, exactly the fault lines your matrix mapped in Chapter 14). "... (Orben & Przybylski, 2019)" foregrounds the finding: use it for consensus sentences where the claim, not

the claimant, is the news. Reviews that are all narrative citations read as book reports (the disease Chapter 16 dissects); all parenthetical reads as evasive. The mixture is the voice.

Citation density signals epistemics. “Social media use rose sharply after 2010” (uncontroversial, one source or none if truly common knowledge); “the association is small and contested (three citations spanning the camps)”; “I argue that...” (no citation: it’s you, and it’s allowed to be you, which is the entire point of Part V). Learning what *doesn’t* need citing, common knowledge in your field, your own argument, your own data, is as much the craft as citing what does.

A citation is a checkable promise: this source, read by me, supports this claim, at this strength. The manager renders the format; only you can supply the promise. Never cite unread, never cite beyond the rung, and never draft with the source’s sentences on the desk.

15.5 The pre-submission citation pass

The mechanics converge in one final ritual, fifteen minutes before any document leaves you (it joins the checklist in Appendix D): every in-text citation resolves to a bibliography entry and vice versa (the plugin guarantees this unless you pasted text between documents: check); spot-verify the five load-bearing citations against their notes: right paper, right claim, right strength (rung audit, one last time); quotes have pages; secondary citations are formatted as such and still deserved; bibliography renders clean (the metadata hygiene payoff, or its overdue invoice); and the style is the assignment’s, not your manager’s default. Fifteen minutes, and the most scrutinized surface of your document, examiners read reference lists the way mechanics listen to engines, is exactly as sound as the system that built it.

15.6 A short gallery of citation errors

The mistakes graders actually see, so you can recognize them in your own drafts:

The et al. overreach: “Verbeek et al.” for a two-author paper, or first use of a three-author work abbreviated when the style wants all names once. The style knows the rule; the plugin applies it; hand-typed citations are where this breeds.

The year mismatch: citing the preprint year for the published version (or vice versa), so the bibliography says 2022 while the text says 2023. Version hygiene (Chapter 13) prevents it; the final read catches it.

The abstract-only citation: citing a paper for a claim that lives only in its abstract, when the body says something more careful. You read at rungs (Chapter 10); cite at rungs.

The moving target: “forthcoming” and “in press” items that published while you drafted. Sweep them at the final pass; their records update in seconds via DOI.

The orphan: a bibliography entry nothing cites, or an in-text citation with no entry, the signature of pasting text between documents. The plugin can’t defend a pasted fragment; the resolution check in the pre-submission pass exists for exactly this.

The Frankenstein quote: a “quotation” assembled from two sentences a page apart, ellipsis-free. Quotes are verbatim or they’re paraphrases; there is no third genre.

Put it to work

- Adopt the standing rule in writing, today, in your project notes: nothing cited without at least pass-1 reading and a note. It will be tested the first time a perfect-looking citation arrives via someone else's paper.
- Practice the plugin workflow for ten minutes: cite three sources in a scratch document, switch the style twice, watch the bibliography follow. The fear dies immediately.
- Audit your notes' QUOTES fields: marks and page numbers present everywhere? Fix now; this is the firewall's maintenance check.
- Take one paragraph you've drafted for any course and run the craft pass: citations moved to their exact claims, narrative-vs-parenthetical chosen deliberately, density matched to controversy. Feel the paragraph sharpen.

Part V

The Literature Review

16

What a Literature Review Is (and Isn't)

Everything in this book has been aimed here. The searches, the passes, the notes, the matrix: all of it exists so that you can write the document this part builds, the literature review, whether it's a standalone assignment, a thesis chapter, or the front third of a paper. And the first job is conceptual, because most weak reviews aren't badly written; they're built to the wrong specification. This chapter fixes the spec.

16.1 The book report, dissected

Here's the opening of a review built to the wrong spec. You've read this review; you may have written it; I certainly did once:

Smith (2018) surveyed 2,000 adolescents and found that heavy social media use correlated with anxiety. Jones (2019) conducted a longitudinal study of 800 teenagers and found small effects. Orben and Przybylski (2019) reanalyzed three large datasets and concluded the associations were tiny. Twenge (2018) argued...

One paper per paragraph (or worse, per sentence), in whatever order the pile happened to sit, each summarized, none connected. It reads like an annotated bibliography that lost its line breaks, and the tell is the grammar: **every subject is an author**. Smith found, Jones found, Orben concluded. The literature is doing all the talking; the writer has vanished into a court stenographer.

Why is this the default failure? Because it's the direct transcription of *serial* reading (Chapter 14 diagnosed this): without a lateral instrument, the pile is the only structure you have, so the pile becomes the outline. The students who wrote book reports weren't lazy; they were matrix-less.

16.2 The actual specification

A literature review is **an argument about the state of a conversation, organized by ideas, in which sources appear as evidence**. Compare the same material to spec:

The field agrees that a statistical association exists between social media use and adolescent anxiety; the live dispute is over its size and direction of cause. Large cross-sectional surveys report moderate correlations (Smith, 2018; Twenge, 2018), but designs with stronger measurement tell a more cautious story: logged-usage and specification-curve studies find effects an order of magnitude smaller (Orben & Przybylski, 2019), and the few longitudinal designs suggest the arrow may partly run from anxiety to use rather

than the reverse (Jones, 2019). The disagreement, this section argues, tracks method rather than population: no study using logged data has found a large effect.

Same five sources. But now the subjects of the sentences are *claims and camps*: the field agrees, the dispute is, the designs tell. The authors have moved into parentheses and evidence roles (the citation craft of Chapter 15, deployed). The paragraph has a thesis, the sources are organized *by their relationship to it*, and the last sentence says something no single source says, a lateral claim, read directly off a matrix fault line (Chapter 14, cashing in). A reader finishes knowing not just what studies exist, but where the conversation stands, and what the writer makes of it.

That's the whole specification, and it's worth stating as the test you'll apply to every paragraph you draft (Chapter 18 operationalizes it): **is the paragraph about an idea, with papers as evidence, or about a paper, with ideas as summary?**

16.3 The jobs a review does

Why do committees, journals, and examiners insist on this document at all? It performs four jobs, and knowing them tells you what yours must contain:

1. **Orientation**: brief the reader on the conversation (Chapter 1's language is now your reader's need): the questions, the camps, the trajectory. Your matrix's consensus columns feed this.
2. **Adjudication**: weigh the evidence where the conversation disputes: which findings are robust, which are method artifacts, what explains the disagreement. Fault-line columns feed this, and it's where critical reading (Chapter 10) becomes visible text.
3. **Gap-making**: establish, honestly, what the conversation hasn't settled, and why that specific gap matters. Empty columns feed this, and the honesty matters: a gap isn't "nobody has done my exact study" (nobody has done anyone's exact study); it's an absence the field itself would recognize as consequential, stated with its reason: unreachable data, new methods only now available, a blind spot the review has just demonstrated.
4. **Positioning**: for thesis and paper reviews, place *your* project as the natural next move in the conversation just described. The review's final gesture: here's the room, here's the open question, here's why my chair is pulled up to exactly that table.

A standalone review assignment weights jobs 1–3; a thesis chapter needs all four, and its examiners will read job 4 as the measure of whether you understand your own project.

16.4 The genre's variants

One spec, several bodies (the genre-awareness of Chapter 1, applied to your own writing):

The narrative review (most course assignments, most thesis chapters): thematic organization, argumentative synthesis, scope declared but search not formalized. This book's default target.

The systematic review (medicine's native form, spreading): the same document with its methods made reproducible: preregistered protocol, logged search (Chapter 8's log, formalized), inclusion criteria, flow diagram, structured extraction (Chapter 14's matrix, formalized). If your field requires it, the workflow you now run is already 80 percent compliant; Appendix C points to the PRISMA standards for the rest.

The mini-review (a paper's introduction or related-work section): jobs 1, 3, and 4 compressed to paragraphs, job 2 mostly implicit. Everything scales down; the idea-organization

rule scales down most importantly of all, because a book-report related-work section is the fastest way to make reviewers doubt a paper.

The thesis chapter: the full four-job version, plus a service function: it teaches your examiners the background *they* need to evaluate your later chapters, which means its scope tracks your project's needs, not the field's whole geography (your scope statement from Chapter 8 was always secretly this chapter's outline fence).

A literature review is not a tour of papers; it's a map of a conversation, drawn by you, with an argument about where the open country lies. Sources are the map's evidence, never its itinerary.

16.5 Permission to have a voice

The book report has one more root worth pulling: students write author-subject sentences because claim-subject sentences feel presumptuous. *Who am I to say "the field agrees"?* Here's the resolution, and it's not a pep talk: **the lateral claims are not opinions; they're findings, and you have the receipts.** "The disagreement tracks method" is a checkable statement about a table you built from notes you can produce, on sources you read critically, gathered by a logged search. That's what the entire workflow was *for*. The voice of a good review is neither the stenographer's (all sources, no writer) nor the pundit's (all writer, no sources); it's the analyst's: claims you own, evidence you cite, and the distance between every "shown" and "said" managed as carefully in your own prose as Chapter 10 taught you to audit it in others'. The next two chapters build that voice a skeleton and then a draft.

16.6 Standing on existing reviews, without being buried

Your corpus almost certainly contains published reviews of your own territory (Chapter 6 sent you to find them), which raises a fair question: if Rev2020 already reviewed this, what is yours for? Three answers, which are also three rules for using reviews well:

Your scope is not their scope. A published review serves its field broadly; yours serves your question (Chapter 4) and, in a thesis, your project's positioning. Their taxonomy is a quarry (Chapter 17 steals section ideas from it); their conclusions are sources to engage, not walls to copy.

Cite the review for the map, primaries for the claims. "For a fuller treatment see Rev2020" is good practice; "passive use predicts anxiety (Rev2020)" when Rev2020 merely summarized Verbeek is citation laundering at one remove (Chapter 15). Load-bearing claims rest on the studies that made them, which you have read.

Newer than the newest review is your natural territory. Every published review has a search-end date, usually visible in its methods. The literature since that date is exactly where your review can add most easily, and "Rev2020's search closed in 2019; the logged-use studies since then change the picture as follows" is about as strong an opening as a student review can have.

Put it to work

- Find one published review in your area (you have one from Chapter 1) and audit five of its paragraphs: subject of each sentence, idea or author? Count the ratio; learn what your field's good reviews sound like.

- Take the worst paragraph of any past review you've written (we all have one) and rewrite it to spec: claim first, sources as evidence, one lateral sentence. Twenty minutes, and the before/after is your personal proof of this chapter.
- Draft the four jobs as one sentence each for your own project: what needs orienting, adjudicating, gap-making, and positioning. These four sentences are the seed of Chapter 17's outline.
- Reread your matrix's fault-line and empty columns and mark the two of each that your review must discuss. You've just chosen your argument's load-bearing walls.

17

From Notes to Structure

Between the full matrix and the first drafted paragraph lies the step students most often skip: building the review's skeleton as its own artifact, on its own day, *before* any prose exists. Skip it and you draft by wandering, discover structure on page six, and rewrite everything (expensively, at deadline). Do it and drafting becomes Chapter 18's calm transcription. This chapter is that day's work, in five moves.

17.1 Move 1: choose the organizing logic

Every review needs one primary axis along which sections follow each other. There are four standard options, and the choice is consequential:

- **Thematic** (the default, and usually right): sections are the major questions or fault lines: measurement, effect size, causality, interventions. Directly supported by your matrix columns; handles disagreement naturally.
- **Chronological**: sections are eras: the early optimism, the replication reckoning, the current synthesis. Right *only* when the conversation's development over time is itself the point; inside sections, time can still order paragraphs.
- **Methodological**: sections are approaches: survey evidence, longitudinal designs, experiments. Right when method is the fault line (our running example flirts with this), or for reviews serving a methods-choosing purpose.
- **Theoretical**: sections are competing frameworks or schools. The humanities' and theory-heavy fields' default.

Hybrids are normal (thematic sections, chronological paragraphs), but the *primary* axis must be singular and chosen on purpose: it's the difference between a review that feels inevitable and one that feels shuffled. What is never an axis: the order you read things in, the alphabet, or the pile.

17.2 Move 2: draft the section map from the matrix

With the axis chosen, sections almost name themselves from the matrix work you flagged in Chapter 16's exercises: each major consensus cluster, each fault line, each consequential gap bids for a section or subsection. Draft the map as a one-page outline: section titles plus, under each, *the claim that section will argue*, in a full sentence. Not "Section 2: Measurement" but "Section 2: How use is measured largely determines what studies find." If you can't write a section's claim-sentence, the section isn't ready (usually it's two sections stapled, or a topic without a point yet: back to the matrix for ten minutes).

For our running example, the map that falls out:

```

1 Intro: scope + why this question, now [scope stmt, Ch7]
2 The association exists, and it is smaller than
  the public debate assumes [consensus col.]
3 Measurement is the fault line: self-report vs
  logged use [fault-line col.]
4 Causality: what panel + experimental designs
  can and can't yet say [fault-line col.]
5 The gap: logged-use longitudinal studies in
  non-WEIRD populations barely exist [empty cols.]
6 (thesis variant) My study: exactly that design

```

Six claim-sentences: the entire review, arguable at a glance. Show this page to an advisor *now*, at the cost of one coffee, and structural feedback arrives while structure is still cheap to change. This artifact is also, not coincidentally, the four jobs of Chapter 16 in document order: orientation (1–2), adjudication (3–4), gap (5), positioning (6).

17.3 Move 3: allocate the evidence

Now walk the matrix once more and **assign every core source to the sections where it will appear** (the note's FOR MY PROJECT field has been predicting this; now it's decided). A spreadsheet column or a tag per section suffices. Three diagnostics fall out mechanically:

- **A source assigned everywhere** is either a genuine anchor (fine, but decide its *home* section, where its full treatment lives; elsewhere it gets one-line cameos) or an unprocessed celebrity you're citing from reflex (Chapter 3 would like a word).
- **A source assigned nowhere** exits the review, however many hours it cost you. This is the sunk-cost moment, and the answer is the same as Chapter 12's mercy: reading that informed your judgment was not wasted by going uncited.
- **A section with two sources** is thin: back to the matrix (are parked papers relevant here?) or back to one targeted search session (Chapter 5, precision mode), or, honestly, the section demotes to a paragraph.

The allocation is also your citation-integrity budget: every assignment is a promise that a note exists to back the coming citation (Chapter 15's standing rule, enforced structurally).

17.4 Move 4: sequence inside sections

Each section is a miniature argument, and its paragraphs need their own order. The pattern that serves almost every section: **claim, strongest supporting evidence, complication or dissent, resolution or assessment**, essentially the paragraph from Chapter 16's specification, scaled up. Concretely, per section of the map: lead paragraph states the section's claim and previews its logic; middle paragraphs each take one sub-idea (never one paper: the spec holds at every scale) with its evidence; a dissent paragraph handles the strongest opposing material *at its strongest* (steelman, Chapter 10, now in your own prose); the closing paragraph assesses: what stands, with what confidence, and hands off to the next section.

Write this as one line per planned paragraph under each section of the map. The one-page outline becomes a two-page skeleton, and notice what you now possess: every paragraph of the review exists as a claim plus assigned evidence, *and no prose has been written*. This is the artifact that makes tomorrow's drafting session start moving in minute one.

17.5 Move 5: the field map (one optional page)

For thesis reviews and any project you'll present orally, one more artifact earns its keep: a single-page visual of the conversation, camps and their relationships, drawn by you: boxes for positions or traditions, arrows for builds-on/contradicts, your project as a marked location. It forces a final synthesis check (if you can't draw it, moves 1–4 have a hole), it becomes the slide every committee meeting wants, and it's the one-glance orientation you'll wish you had at the viva. Thirty minutes, pencil first.

Structure is a decision, made once, on evidence, before drafting, not a discovery made on page six. Axis, map, allocation, sequence: when every paragraph exists as a claim with assigned sources, the hard intellectual work of the review is finished, and it happened without a single sentence of prose.

17.6 A word on the blank-page terror

Notice what this chapter did to the mythology of writing. The terror of the blank literature review assumes the page is where thinking happens: you sit down, sixty papers swirling, and must conjure order live. Everything in this book has been dismantling that assumption piece by piece: the notes think per-paper (Chapter 11), the matrix thinks laterally (Chapter 14), this chapter's moves think structurally, and each stage's artifact carries the thinking forward. By the time prose begins, every hard question has been answered somewhere you can point to. The blank page was never the problem; arriving at it empty-handed was. You're not arriving empty-handed.

17.7 The advisor in the loop

Structure is the cheapest stage for feedback, and the easiest to ask for badly. Two rules make it land:

Match the artifact to the question. The one-page map asks a strategic question: *are these the right sections, is this the right axis, is the gap credible?* The two-page skeleton asks a tactical one: *is the evidence allocated sanely, is anything thin?* Sending a full draft to ask a strategic question wastes both of you; structural feedback on page six of prose arrives too late to be cheap.

Translate the feedback you'll get. "This reads like a list" means the claim-sentences are missing or hidden: re-run the section-map move. "Where are you in this?" means assessments and weighing sentences are absent: the synthesis polish of Chapter 18 is the cure. "What about the X literature?" is a new column and a targeted search, not a crisis. And "this is too broad" is your scope statement (Chapter 8) asking to be re-read, then obeyed.

One special case: repurposing a course review for the thesis. It's legitimate scholarship practice (with disclosure per your institution's rules, Chapter 15), but it's a re-scoping job, not a paste: the thesis review serves your project's positioning, so re-run the allocation against the new scope and expect a third of the material to leave.

Put it to work

- Choose your axis, in writing, with one sentence on why (and one on the runner-up you rejected; the comparison sharpens the choice).

- Build the one-page section map with claim-sentences. Time-box: ninety minutes. Send it to your advisor or a peer the same day.
- Do the allocation pass: every core source assigned or consciously cut; every section's source-count checked. Note your thinnest section and decide its fate (reinforce, demote, or targeted search).
- Expand to the two-page skeleton: one line per paragraph, claim plus citekeys. Schedule Chapter 18's first drafting session for the day after it's done, while the skeleton is hot.

18

Drafting, Revising, and Keeping It Alive

The skeleton is built, the evidence allocated, every paragraph a claim with citekeys attached. What remains is prose, then revision, then the discipline of keeping the review current while the rest of your project unfolds. This is the longest chapter in the book because it's the last mile, and the last mile is where reviews are won: two drafts of the same skeleton can read as authority or as homework, and the differences are specific, learnable moves. Here they are.

18.1 Drafting from the skeleton

The session mechanics first. Draft one section per sitting, from the two-page skeleton, **with your notes open and the PDFs closed** (Chapter 15's patchwriting firewall, now operational). Each paragraph line in the skeleton becomes a paragraph of prose: state the claim, deploy the evidence from the notes' own-words CLAIM fields, cite through the picker as you go (never "I'll add citations later," which is how orphan claims and midnight bibliography sessions are born). Where a cell of the matrix carries the comparison, write the lateral sentence directly from it.

Draft *ugly on purpose*. The session's goal is a complete section with every claim cited, not beautiful sentences; polish is a separate pass (below) and doing both at once does both badly. A skeleton section drafts in sixty to ninety minutes this way, which means **a full first draft of a thesis-chapter review is roughly a week of sessions**, a claim students disbelieve until the skeleton makes it true.

Two localized crafts while drafting:

Paragraph openers carry the argument. Read your draft's first sentences alone, in order: they should tell the review's whole story (claim, claim, complication, claim, assessment). If an opener is "Another study..." or "Smith (2018) found...", the paragraph has lost its spec (Chapter 16); fix the opener first and the paragraph usually follows.

Attribution verbs are calibrated instruments. "Shows" claims more than "finds," which claims more than "suggests," which differs from "argues" (a position, not a result) and "claims" (audible skepticism). Match the verb to the rung you actually bought in your critical read (Chapter 10): a paper whose evidence you rated as suggestive gets "suggests," whatever its abstract said. Tense, while we're here: your field has a dialect (APA-style fields: past tense for findings, "Smith found"; many humanities: present, "Smith argues"); copy your model review (Chapter 16's audit) and be consistent.

18.2 The synthesis polish

Draft complete, the first revision pass targets the one quality that separates authority from homework: density of *connection*. Three moves, applied paragraph by paragraph:

Convert sequence into relation. Wherever two sourced claims sit adjacent without a relation word, decide the relation and say it: *likewise*, *by contrast*, *extending this*, *complicating this picture*. The relation words are your matrix links (Chapter 11's Connects-To) surfacing in prose, and they're the sinew a book report lacks.

Lift summaries into assessments. Find paragraphs that end on a source ("...with effects persisting at follow-up (Jones, 2019)") and add the writer's move: what does the section now know? ("Together these designs make simple exposure explanations hard to sustain.") Every paragraph earns its assessment sentence or confesses it was inventory.

Weigh, don't count. Hunt the phrase "many studies" and its cousins: replace counting with weighing. Not "many studies find effects" but "the largest and best-instrumented studies converge on small effects; the large effects appear only where measurement is weakest," if that's what your matrix shows. Weighing sentences are the review's most valuable, and they're only writable from a matrix, which is why homework reviews can't fake them.

18.3 The honesty audit

Second pass, shorter, sharper: the review's integrity surfaces. Check the dissent is at full strength where you argued against it (would its author sign your summary of them? The steelman test, in public now). Check your gap survives adversarial reading: is it truly unaddressed within your logged scope, and would the field care? Check the hedges: every "may," "might," "appears to" either earned by genuine uncertainty or deleted; hedging as politeness reads as not knowing your own material. And check the confidence: every flat assertion attached to either a citation, your matrix, or your explicit argument. The two failure styles, timid and grandiose, are both calibration failures, and this pass is where calibration becomes a property of the text (Chapter 10's discipline, turned on yourself, one last time).

Then the mechanical closers, borrowed whole from the companion craft: the citation pass of Chapter 15, and, if your document lives in LaTeX, the build hygiene of my thesis book. The full pre-submission sequence is Appendix D.

18.4 The living review

For any project longer than a semester, the review isn't finished when drafted; it's *versioned*. The conversation continues (your alerts from Chapter 12 are the ears), your project itself will shift what matters, and the document must keep pace without consuming you. The maintenance contract, three clauses:

New arrivals join through the system, not the text. An alert surfaces a relevant paper: it gets the filing ritual, the queue, a pass, a note, a matrix row, *then* a decision about the text: does it slot into an existing paragraph (one sentence, one citation), strengthen or wound an existing claim (revise that claim honestly, including your gap if the gap just got addressed: better you find that paper than your examiner), or change nothing (most common: the note exists, the text stands). Never paste a new source directly into prose ahead of the pipeline; that's how unread citations breed.

Reconcile at milestones. At each project milestone (proposal, mid-point, pre-submission), one scheduled session re-reads the review against the current matrix: claims still supported? Confidence still calibrated? The scope statement and search log get their

final update here too, because the methods paragraph they feed (Chapter 8) ships with the thesis.

Freeze deliberately. In the final weeks, declare the review closed except for genuinely load-bearing arrivals (new results in your exact niche: yes; another replication of the consensus: logged, not written). An endless-update review is the search-that-never-stops (Chapter 8) in a new costume, and the same permission applies: convergence, documented, is the standard; omniscience is not.

Draft from notes with the sources closed; revise for connection, then for honesty; and afterward, let new literature reach the text only through the pipeline. The review is the workflow's visible surface, and it stays exactly as healthy as the system underneath it.

18.5 Session logistics

Small mechanics that make drafting sessions start fast and end well:

The desk setup: skeleton on one side, notes (the manager) on the other, the draft in the middle, PDFs *closed* (the firewall), phone elsewhere. The citation picker hotkey learned (Chapter 13), so citing doesn't break the sentence.

The warm start: begin each session by re-reading only the previous session's last paragraph and the skeleton line you're about to write, not the whole draft. Whole-draft re-reading at session start is the most respectable form of not writing.

The rough-in rule: when a sentence fights you, type the ugly version in brackets ([something about how attrition weakens this]) and keep moving; the polish pass exists precisely so drafting doesn't have to be good, only complete.

The downhill stop: end sessions mid-section, one paragraph *before* you're empty, and leave a one-line note about what comes next. Tomorrow's session then starts downhill. Writers have used this trick for a century because it works.

The session log: one line per session (date, section, word count, next step) in the same file as the skeleton. Momentum you can see is momentum that survives a bad week.

Put it to work

- Draft your first section tomorrow, from the skeleton, notes open, PDFs closed, sixty to ninety minutes, ugly on purpose. Then draft the rest at one section per session.
- Run the synthesis polish on your first drafted section the day after: relation words in, assessment sentences in, counting converted to weighing. Note how the same facts change register.
- Run the honesty audit with a peer swap if you can trade (Chapter 10's fairness rules apply to feedback in both directions).
- Put the milestone reconciliations in your calendar now, and write the freeze date next to your deadline. The review is alive from today, and, from today, under control.

19

Review Makeovers

Theory is done; this closing chapter is an operating room. Four passages of the kind that appear in real student reviews, each diagnosed and rebuilt using the tools this book installed. Read them slowly, because these four patients cover the failures that account for most lost marks on literature reviews, and the rebuilds show the whole workflow surfacing in prose. All four continue the running example, so you can watch the same corpus you’ve followed since Chapter 4 being handled badly and then well.

19.1 Makeover 1: the book report

BEFORE

Twenge (2018) analyzed large-scale survey data and found that adolescents who spent more time on screens were less happy and more anxious. Orben and Przybylski (2019) used specification curve analysis on three datasets and found very small effects. Jones (2019) conducted a two-wave panel study and found small reciprocal associations. Verbeek et al. (2023) used experience sampling with logged data and found that passive use predicted next-day anxiety.

Diagnosis. Chapter 16’s specimen disease in the wild: every sentence’s subject is an author; the order is arbitrary (it’s actually the order read, the pile as outline); no sentence relates any study to any other; and the paragraph ends without having claimed anything. A grader learns that the student read four papers. Nothing else.

AFTER

The association between social media use and adolescent anxiety is real but easily overstated, and how it is measured largely determines what is found. Studies relying on self-reported use report moderate correlations (Twenge, 2018), but designs with stronger measurement consistently shrink the effect: reanalyses controlling for analytic flexibility find associations an order of magnitude smaller (Orben & Przybylski, 2019), and the only studies using logged rather than self-reported use find small, short-term, within-person effects confined to passive scrolling (Verbeek et al., 2023). Panel evidence further complicates the causal story, with associations running in both directions (Jones, 2019). No study using logged data has yet found a large effect.

What changed. The paragraph now opens with a claim (the section-map claim-sentence of Chapter 17), the sources enter as evidence for positions rather than as protagonists (Chapter 15’s parenthetical choice, made deliberately), relation words carry the argument (“but,” “further complicates”), and the closing sentence is a lateral claim read

straight off the synthesis matrix (Chapter 14): a sentence no single source contains, which only the writer of this paragraph could have written. That sentence is the difference between reporting a pile and knowing a field.

19.2 Makeover 2: the citation carpet

BEFORE

Many studies have found that social media use is associated with anxiety in adolescents (Twenge, 2018; Smith, 2018; Jones, 2019; Verbeek et al., 2023; Kim et al., 2022; Rees, 2021). However, some studies disagree (Orben & Przybylski, 2019).

Diagnosis. Counting instead of weighing (Chapter 18's third polish move, violated), and worse, the count is dishonest: the six carpeted citations include a qualitative study (Rees) that measured no association, and the studies "agreeing" disagree with each other about size by an order of magnitude. The carpet also buries the actual structure of the disagreement, reducing the field's most interesting finding, that method predicts result, to "however, some disagree." Six citations, and the reader knows less than after Makeover 1's rebuild.

AFTER

The heaviest-measured studies find the smallest effects. Large self-report surveys report moderate associations (Twenge, 2018; Smith, 2018), while the logged-use designs, though fewer, converge on small within-person effects specific to passive use (Verbeek et al., 2023; Kim et al., 2022), and specification-curve reanalysis suggests much of the remaining variation reflects analytic choices rather than the phenomenon (Orben & Przybylski, 2019). Qualitative work points to mechanisms, social comparison during passive scrolling, that are consistent with this pattern (Rees, 2021). Weighing design quality rather than counting studies, the evidence favors a small, mechanism-specific effect over the large harms of public discourse.

What changed. Citations moved from carpet to claims, each attached to exactly what it supports (Chapter 15); the qualitative study got the mechanism role it actually plays (Chapter 10's dialects); and "many studies" became a weighing argument with its principle stated. Note that the rebuild has the same six sources; not one new paper was needed, only the matrix's sight lines.

19.3 Makeover 3: the fake gap

BEFORE

While much research has examined social media and anxiety, no study has specifically examined the relationship between passive Instagram use and anxiety among 14-to-16-year-olds in Tier-2 Indian cities using experience sampling. This study addresses this important gap.

Diagnosis. The "nobody has done my exact study" gap, and Chapter 16 warned it convinces no one: intersect enough qualifiers and every study on earth is unprecedented. The passage never argues the gap *matters*, never shows the field would recognize it, and "important" is asserted where it needed to be earned. An examiner reads this as a student who found a hole in a bibliography rather than a hole in knowledge.

AFTER

Two limits of the current evidence motivate this study. First, the logged-use designs that anchor the field's most credible estimates (Verbeek et al., 2023; Kim et al., 2022) have been conducted entirely in European and East Asian samples; whether their small within-person effects generalize to contexts with different platform mixes and usage norms is untested, and reviews have flagged the WEIRD concentration of this literature as a priority (Rev2020). Second, no logged-use study has yet run in a setting where data costs meaningfully constrain passive scrolling, a condition common across South Asia that plausibly moderates the passive-anxiety mechanism itself. An experience-sampling design in Tier-2 Indian cities addresses both: it tests the field's core finding out of sample, where existing theory gives reasons to expect the effect to differ.

What changed. The gap became an argument with two legs: the field's own map shows the absence (empty matrix columns, cited to the review that flagged it), and a *reason* the absence matters (the effect could plausibly differ, and here's the mechanism why). The qualifier pile-up became a designed study answering a recognizable open question. Same project; the difference is entirely in whether the gap was found in the literature or argued from it.

19.4 Makeover 4: the hedge fog**BEFORE**

It could be argued that measurement may play a role in the differences between studies. It seems that logged-use designs might find somewhat smaller effects, which could suggest that self-report methods may possibly overestimate associations, although this may not always be the case.

Diagnosis. Nine hedges in two sentences, and not one earned: the claim being smothered, that measurement predicts effect size, is one of the best-supported patterns in the writer's own matrix. This is the timid calibration failure of Chapter 18's honesty audit: hedging as politeness, which reads as not knowing your own evidence. (The grandiose twin, stripping all hedges from genuinely uncertain claims, fails the same audit in the other direction.)

AFTER

Measurement method is the strongest predictor of effect size in this literature. Every logged-use study in this review finds smaller associations than every self-report study, a pattern consistent with documented gaps between recalled and actual use (Smith, 2018). Whether the remaining small effect is causal is genuinely uncertain: the panel evidence is bidirectional (Jones, 2019), and no experimental manipulation of passive use has yet been attempted with adolescents.

What changed. The supported claim now stands flat, with its evidence ("every... than every" is a matrix sentence, checkable against the table); the hedges migrated to where uncertainty actually lives, the causal question, and became *specific*: not "may not always be the case" but exactly what's missing (no experiments) and what's ambiguous (direction). Calibration, in prose: confidence proportional to evidence, doubt located precisely.

19.5 The pattern across all four

Read the diagnoses again: authors as subjects, counting as argument, gaps by qualifier-stacking, hedging as reflex. All four are what prose looks like when it's generated from a *pile*. And every rebuild drew on the same three artifacts: the matrix (for lateral claims and weighing), the notes (for rungs, mechanisms, and honest limits), and the section map (for claims to open paragraphs with). Not one rebuild required new reading. **Weak reviews are almost never under-read; they're under-systemed**, which is this book's thesis, demonstrated four times on one page each.

19.6 The whole workflow, in one page

So here it is, the complete system this book promised, at postcard scale:

1. **Orient** (Part I): know the genres, the anatomy, and the trust calculus of the conversation, and distill your topic into a typed, piloted, written question.
2. **Find** (Part II): concept-table searches and citation-web walks, papers captured through the access ladder into the manager, everything logged, stopping when the log converges.
3. **Read** (Part III): three passes allocating attention, critical reading measuring shown against said, structured notes banking it all, a queue pacing it weekly.
4. **Organize** (Part IV): one manager as source of truth, the matrix as lateral instrument, citations as checkable promises.
5. **Write** (Part V): a review specified as an argument, structured before drafting, drafted from notes, revised for connection and honesty, kept alive by the same pipeline, and, when it falters, repaired with the four makeovers above.

Every stage leaves artifacts; every artifact feeds the next stage; nothing depends on heroics, memory, or mood. That's what a workflow is, and it's yours now. The appendices hold the templates and checklists; the conversation holds a seat with your name on it. Go say something.

Put it to work

- Run the four diagnoses over your own current draft: mark every author-subject paragraph opener, every “many studies,” every qualifier-stacked gap, every hedge. Just mark, today.
- Rebuild your worst paragraph using its matrix row and notes, before/after style. Keep the pair; it's your personal proof of the method, and a startlingly effective thing to show a struggling classmate.
- Check your gap paragraph against Makeover 3's two-leg standard: absence shown from the field's map, and a reason the absence matters. If either leg is missing, you know the fix.
- When the review ships, run Appendix D one last time, then close the laptop. You've joined the conversation properly, and the next project inherits everything: the system, the library, and a researcher who knows how to use them.

Part VI

Toolkit

A

The Paper Note Template

The template from Chapter 11, ready to copy into your reference manager as a paste-snippet, followed by a worked example at full pass-2 depth, and the pass-1 short form. A printable companion lives at gauravtiwari.org.

The template

```
CITEKEY:                                STATUS: pass N (date)
IN ONE LINE:
CLAIM(S):      (licensed findings, my words, magnitudes)
DESIGN:        (who/what, n, method, comparison)
EVIDENCE NOTES: (key results; figures worth returning to)
LIMITS THAT MATTER: (2-3, mine as well as theirs)
STEELMAN:      (their best case, one sentence)
QUOTES:        (" verbatim + page numbers, sparse)
MY REACTION:   (what this does to my picture; "?" list)
CONNECTS TO:   ([citekeys] agrees/contradicts/extends/method)
FOR MY PROJECT: (theme, section, role)
```

Worked example, pass-2 depth

```
CITEKEY: verbeek2023passive    STATUS: pass 2 (2026-06-18)

IN ONE LINE: 8-week experience-sampling study arguing
passive scrolling, not posting, predicts next-day anxiety
in teens.

CLAIM(S): (1) Within-person: higher passive-use days predict
small next-day anxiety rise (b = 0.11 on GAD-7-adapted
scale); active use shows no association. (2) Between-person
differences mostly vanish once controls added. Licensed
rung: short-term within-person association, NOT causation,
NOT long-term.

DESIGN: experience sampling, n = 384 (ages 13-17, two Dutch
schools), 5 prompts/day x 8 weeks, app-logged use split
passive/active, multilevel models.

EVIDENCE NOTES: Fig 2 = the whole story (within-person
```

slopes); Table 3 sensitivity checks hold. Compliance 71%, reported honestly.

LIMITS THAT MATTER: two schools, one country; anxiety scale adapted (validation thin, my flag not theirs); passive/active split via app category is crude (their flag).

STEELMAN: even with crude splits, logged within-person design beats nearly all prior self-report work on this exact question.

QUOTES: "passive use may function as both symptom and amplifier" (p. 1108)

MY REACTION: strongest design in my pile on the passive question; moves me toward "small but real" camp. ? = does anyone replicate outside Europe? ? = is symptom-vs-amplifier testable at all in 8 weeks?

CONNECTS TO: extends orben2019screens (answers its measurement complaint); contradicts twenge2018 magnitude; method-cousin of jensen2021ema.

FOR MY PROJECT: core. Themes: passive-active, measurement. Home section: 3 (measurement fault line); cameo in 4 (causality).

Fifteen minutes of typing at pass-2 speed. Note what makes it survive: every field self-contained, magnitudes present, the rung stated, quotes firewalled with a page number, the "?"'s honest, and the last two fields already doing Part IV and Part V's work.

The pass-1 short form

CITEKEY: kim2022replication STATUS: pass 1 (2026-06-20)
 IN ONE LINE: Korean replication of verbeek-style ESM design, smaller n, same direction.
 DECISION: park -- replication, adds non-EU sample; revisit if section 3 needs non-Western evidence.

Two minutes, and the decision line is the part you'll thank yourself for: it's the difference between re-triaging this paper in month four and remembering, instantly, why it waits where it does.

B

The Synthesis Matrix, Worked

Chapter 14’s instrument, shown small. A slice of the running example’s matrix, six sources by four themes (a real one would run 15–30 rows and five to eight columns; the mechanics don’t change):

Source	Effect found	Measurement	Causality	Use type
twenge2018	moderate assoc.	self-report, cross-sect.	implies use→harm	—
orben2019-screens	tiny ($r \approx .05$)	self-report, spec-curve	agnostic; attacks overclaim	—
jones2019	small, bidirectional	self-report, 3-wave panel	partly anxiety→use	—
verbeek2023-passive	small, within-person	logged , ESM	assoc. only, short-term	passive only
kim2022-replication	small, same dir.	logged, ESM	assoc. only	passive only
rees2021qual	—	interviews	mechanism hypotheses	supports split

Read the columns the way Chapter 14 taught:

- **Effect found** shows the fault line, and read against **Measurement**, shows it tracking method: the moderate effect sits in self-report cross-section; every logged design reports small. That correlation is the review’s central lateral claim, visible in one glance.
- **Causality** shows a column mostly agnostic or bidirectional, which converts “social media causes anxiety” from a premise into an open question your review must treat as one.
- **Passive/active** is two-thirds dashes: an empty-column gap, and note *who* fills it (only the ESM studies), which tells you the gap is instrumentation-shaped, worth saying in exactly those words.

Building yours: the mechanics

1. One spreadsheet, first sheet the matrix, second sheet your theme-word list and scope statement (everything in one place travels well to advisor meetings).
2. Column A citekeys, B year, C the DESIGN one-liner, then theme columns. Freeze panes at column D.
3. Cells: one to two compressed lines, magnitudes kept, dashes for silence, and **never** content your note can’t back (the matrix is a view of the notes, not a second database).

4. Row fill rides the pass-2 ritual (five minutes at note time); column additions get a dedicated retrofit half-hour.
5. Sort experiments are free insight: sort by year to see the conversation move; by design to see the method clusters; by any theme column to draft that section's paragraph order (Chapter 17, move 4, nearly automates itself).

A blank starter file with these conventions is at gauravtiwari.org.

C

A Field Guide to Sources and Databases

Chapter 5’s advice, made concrete per field: where the conversation lives, which database gives the systematic pass, and the local quirks worth knowing. Names not links (interfaces move; your library’s website has the current doors), and your subject librarian remains the living edition of this appendix.

Cross-field, everyone

Google Scholar (exploration, cited-by, alerts), **Scopus** and **Web of Science** (curated multidisciplinary indexes, strong citation tooling, the usual home of “systematic” searches outside medicine), **your library discovery layer** (one search across everything subscribed), **ProQuest Dissertations** (theses, Chapter 1’s underrated genre), and **Retraction Watch’s database** (the what-happened-next check of Chapter 3).

By field

Medicine and life sciences. PubMed/MEDLINE with MeSH controlled vocabulary (the model all controlled vocabularies aspire to); Embase for the systematic-review second database; Cochrane Library for existing systematic reviews (always check before starting one); bioRxiv/medRxiv preprints, weighted per Chapter 3’s caution; PRISMA guidelines when your review must be systematic.

Psychology. PsycINFO (APA thesaurus = the controlled vocabulary); PubMed overlaps the clinical side; the replication-era caveats of Chapter 10 apply at full strength here, so weight preregistered and replicated work explicitly in your matrix.

Computer science and engineering. Conferences are the conversation (Chapter 1); DBLP for clean CS bibliography, IEEE Xplore and the ACM Digital Library for full text, arXiv as the field’s bloodstream (with versions noted); Semantic Scholar’s tooling is strong here. Engineering proper adds Inspec; Compendex.

Mathematics and physics. arXiv first; MathSciNet and zbMATH (reviews of papers, a genre bonus: an expert’s one-paragraph assessment of many papers, free critical reading); INSPIRE-HEP for high-energy physics.

Social sciences. Scopus/WoS as the spine; EconLit (economics, with its working-paper culture: NBER and SSRN predate “preprints” by decades); ERIC for education; Sociological Abstracts; national statistical agencies and the gray literature of Chapter 3 carry real weight in applied questions.

Humanities. MLA International Bibliography (literature, languages), JSTOR and Project MUSE (deep backfiles, monograph chapters), L'Année philologique (classics), RILM (music), field-specific bibliographies compiled by scholarly societies; the monograph and the edited-volume chapter are first-class citizens (Chapter 7's ILL rungs get heavy use), and book reviews in journals are a critical-reading shortcut the sciences lack.

Law. Westlaw/LexisNexis (institutional), SSRN's legal studies network, HeinOnline for history; citation practice is its own world (your style guide is probably the Bluebook or OSCOLA, and Chapter 15's principles hold while the mechanics differ).

Preprint servers, by neighborhood

arXiv (math, physics, CS, stats), bioRxiv/medRxiv (life, medical), SSRN (social science, law, economics), PsyArXiv, SocArXiv, ChemRxiv, EarthArXiv, and institutional repositories everywhere (the green-OA rung of Chapter 7). Weight per Chapter 3: current, unrefereed, check for the version of record before citing.

Choosing your two

The working setup for any single project is smaller than this appendix: **Scholar plus one field database** covers the exploration/systematic pair of Chapter 5 for almost everyone. Pick the field database by asking one question of an advisor or librarian: "if you could only search one database for this topic, which?" Then log both (Chapter 8), alert both (Chapter 12), and stop collecting databases the way this book taught you to stop collecting papers.

D

The Literature Review Checklist

Run before any review leaves your hands: assignment, proposal chapter, thesis chapter, or a paper's related work. Each item names its chapter.

The search behind it (Part II)

1. Scope statement written, current, and honored, exclusions included (Ch. 8).
2. Search log shows systematic passes in at least two tools and convergence toward zero new core finds (Chs. 5, 8).
3. Foundations identified via snowballing, and their afterlives checked (Chs. 6, 3).
4. Alerts running; a dated freeze decision exists (Chs. 12, 18).

The reading behind it (Part III)

5. Every cited source has at least a pass-1 note; every load-bearing source a pass-2/3 note (Chs. 9, 11).
6. Claims recorded and cited at their licensed rung, with magnitudes where they matter (Ch. 10).
7. The strongest dissent is present, at steelman strength (Chs. 10, 18).

The document itself (Part V)

8. Organizing axis chosen deliberately; sections carry claim-sentences, not topics (Ch. 17).
9. Paragraphs pass the spec test: about ideas, sources as evidence; paragraph openers alone tell the argument (Chs. 16, 18).
10. Lateral claims present (consensus, fault lines, weighing sentences) and each backed by the matrix (Chs. 14, 18).
11. The gap is honest: within scope, checked against the log, and stated with why it matters (Chs. 8, 16).
12. (Thesis/paper) Positioning closes the review: your project as the conversation's next move (Ch. 16).
13. Attribution verbs calibrated; hedges earned; tense consistent with your field's dialect (Ch. 18).

Citations and mechanics (Part IV)

14. Nothing cited unread; secondary citations formatted as such and still deserved (Ch. 15).
15. Quotes marked with pages; no drafting happened from open PDFs (patchwriting firewall held) (Chs. 11, 15).
16. Citations attached to their exact claims; narrative/parenthetical chosen deliberately (Ch. 15).
17. Every in-text citation resolves to the bibliography and vice versa; five load-bearing citations spot-verified against notes; style is the assignment's (Chs. 13, 15).
18. The bibliography read once, end to end, for metadata casualties (Ch. 13).

Eighteen items; with the workflow running, most are already true and the pass takes half an hour. If more than a few fail, the review isn't underpolished, it's undersystemed, and the fix is the chapter named beside the failure, not a longer night.

E

Tools, Honestly

The current inventory, with the biases declared: this book prefers free, durable, export-friendly tools, and treats every subscription as a question of what happens to your library when you stop paying. Categories in workflow order; product details drift, so the *criteria* in each section are the durable content.

Reference managers (Ch. 13)

Zotero: the book's default. Free, open-source, nonprofit; best-in-class capture; PDFs, notes, tags, groups; plugin ecosystem (Better BibTeX for LaTeX lives). Weaknesses: storage above the free tier costs (modestly), interface is workmanlike. **JabRef**: open-source, .bib-native; right for pure-LaTeX workflows that want the bibliography file as the primary artifact. **Paperpile**: polished, subscription, Google-Docs-native; fine if you live there; check export before committing. **EndNote**: paid, heavyweight, institutional; adopt only where a lab standardizes on it. **Mendeley**: hard to recommend new adoptions (ownership incentives, feature retreat). The criteria that outlive this paragraph: one-click capture, real export, citation-key control, word-processor integration, and a survivable answer to “what if this product dies?”

Search and citation-graph tools (Chs. 5, 6)

Beyond Scholar and the library databases (Appendix C): **Semantic Scholar** (clean meta-data, influence-weighted citations, good API); **Connected Papers**, **Research Rabbit**, **Litmaps** (visual citation neighborhoods; use per Chapter 6: seed, explore, harvest with a why-note, never browse-and-forget); **Unpaywall** and kin (legal OA detection, Chapter 7's rung 2). Criteria: does it show *why* items are related, does it export to your manager, does it index your field's actual venues?

Note systems (Chs. 11, 14)

The book's default is deliberately boring: **notes in the reference manager** (anchored to items, synced, searchable, zero extra tools) plus **a spreadsheet matrix**. The step-up, for multi-year intellectual lives: **Obsidian** (local Markdown, links, plugins; Zotero-integration plugins mature), **Logseq** (similar, outline-first), both durable-by-format. The criteria: own-your-files, links as first-class content, citekey compatibility, and honestly, whether the tool serves the notes or the notes serve the tool; a note system you tinker with weekly is a hobby, not an instrument.

AI tools, with both eyes open

The category moves fast (this is 2026's snapshot; the criteria will outlive it). Current tools summarize papers, answer questions “from” documents, suggest related work, and draft prose. Where they genuinely help a student workflow: **triage acceleration** (a generated summary as a pass-1 *aid*, never replacing your own decision line); **comprehension support** (explaining an unfamiliar method at pass 2, like an infinitely patient labmate whose claims you verify); **search vocabulary** (suggesting the field's terms for your concept table).

Where the workflow must not delegate, and why:

- **Reading of record.** An AI summary is a secondary source of unknown reliability: models compress confidently, hallucinate citations, and miss the exact rung-level nuance Chapter 10 taught you to extract. Your notes must come from your reading; the standing rule of Chapter 15 (cite nothing unread) has no AI exception.
- **The lateral claims.** Consensus, fault lines, gaps: these must trace to *your* matrix over *your* corpus, or your review's central sentences are unverifiable hearsay.
- **Drafting from model output.** Beyond integrity policies (check your institution's, and disclose per Chapter 15's honesty norms), model prose is patchwriting-shaped by construction: fluent, sourceless, and calibrated to sound right rather than be right. Draft from your notes; use tools, if at all, at the polish margins your policy permits.

One sentence of standing advice: **use AI where errors are cheap and verification is instant; never where the workflow's integrity chain (read → note → cite) would break silently.**

Writing stacks (Part V, and the companions)

Word/Google Docs + manager plugin: the default; everything in Chapter 15 works today. **LaTeX + biblatex:** the scientific long-form stack; my *LaTeX for Theses and Dissertations* is that book, and Better BibTeX bridges your Zotero library into it with the same citekeys your notes already use. **Typst and kin:** promising newer typesetters; check citation-style coverage against your requirements before betting a thesis on one.

F

Further Reading

Short and honest, each entry with when it's worth your time.

S. Keshav, "How to Read a Paper." The two-page classic behind Chapter 9, free online. Read it once for the original; hand it to friends who won't read a book.

Umberto Eco, *How to Write a Thesis*. Written for 1977's index cards, and the thinking underneath, scoping, note discipline, the ethics of sources, remains the most humane treatment ever published. Read for spirit, not tooling; this book is, in part, Eco's workflow rebuilt for your century.

Diana Ridley, *The Literature Review: A Step-by-Step Guide for Students*. The standard academic-press treatment of Part V's territory, with many worked excerpts from real theses. Useful when you want a second voice on structure and more example passages than my Chapters 16–18 had room for.

Booth, Colomb, and Williams, *The Craft of Research*. The classic on the whole arc from question to argument. Broader and more rhetorical than this book's workflow focus; the chapters on making claims and acknowledging objections deepen Chapters 10 and 18 well.

Sönke Ahrens, *How to Take Smart Notes*. The book behind the linked-note movement (Chapter 14's step-up path). Read between projects, not during one; adopt its principles (own words, links, one idea per note), and be warned that its tooling enthusiasm has consumed whole semesters.

The PRISMA statement (and your field's reporting guidelines). When your review must be systematic, these are the specification; free online, updated periodically, and your methods section's skeleton.

Ben Goldacre, *Bad Science*. The most enjoyable route into Chapter 10's instincts: evidence pyramids, publication bias, and statistical self-defense, taught through gleeful demolition. A vaccine against credulity that a first-year can read on a weekend.

And from the same shelf as this book: *How to Write Mathematics* for the writing craft in quantitative fields, and *LaTeX for Theses and Dissertations* for the document your research will finally live in. The three are one workflow, told in thirds.

G

The Emergency Protocol

The honest appendix. The review is due in about a week, the workflow was not in place, and you're holding this book like a fire extinguisher. Here is the whole system, compressed to seven days. It produces a legitimate, defensible review, smaller in scope than the full-workflow version, and it works precisely because it *is* the workflow, run at emergency compression rather than abandoned. Chapters cited so you can go deeper where you have an hour to spare.

Day 1: question, scope, and the map (3 hours)

Run the funnel (Chapter 4) fast: topic to question in forty minutes, on paper. Write a five-line scope statement with aggressive exclusions (Chapter 8); in emergency mode, scope is your best friend, and narrow is survivable while broad is not. Then spend the rest of the session hunting **one or two recent review articles** (Chapter 5: topic + “review”): a good review is a hundred hours of someone else's work, and this week, you're allowed to stand on it openly.

Day 2: search and harvest (3 hours)

Build the mini concept table (Chapter 5), run one broad Scholar pass and one recent-years pass, then backward-snowball the review's references and forward-snowball its citers (Chapter 6). Harvest into a reference manager, and yes, install Zotero today (Chapter 13: the forty-five-minute setup pays even in week one, because the bibliography will generate itself on day 7, which is exactly when you won't have time to hand-format one). Target: 20–30 candidates, filed, with one-line why-notes. Log your searches in ten lines (Chapter 8); even an emergency review may need a “how I searched” paragraph.

Day 3: triage and core selection (2–3 hours)

Pass-1 the entire harvest (Chapter 9): five minutes each, decision line each. Select **six to ten core papers**, no more; an emergency review argued from eight well-read sources beats a skimmed twenty-five, and examiners can tell which they're reading. Everything else is parked with its one-liner (some will surface as single-sentence citations later; that's their whole role this week).

Days 4–5: read and note the core (2 sessions each)

Pass-2 each core paper (Chapters 9, 10): an hour each, figures first, the template note at working depth (Chapter 11, Appendix A), rungs respected, magnitudes captured. At the end of day 5, build the **mini matrix** (Chapter 14, Appendix B): your core rows, four or five theme columns, forty-five minutes. Read it and write down, in sentences, one consensus, one fault line, one gap. Those three sentences are your review's spine.

Day 6: structure, then draft (the long day)

Morning: the section map with claim-sentences (Chapter 17), evidence allocated from the matrix; keep it to four or five sections. Afternoon and evening: draft from notes with PDFs closed (Chapter 18), one section per sitting, ugly on purpose, citing through the picker as you go. A 2,500-word review drafts in three sittings from a real skeleton; it cannot be drafted at all from a pile, which is why days 3–5 were not skippable.

Day 7: polish, audit, ship (3 hours)

The synthesis polish (relations in, counting to weighing), the honesty audit (steelmanned dissent, earned hedges, an honest gap), then the citation pass and the checklist (Chapters 18, 15, Appendix D). Read the paragraph openers top to bottom once: they should tell the argument alone (Chapter 18). Ship it.

What you sacrificed, and what you didn't

Sacrificed, and worth saying so in the review's scope sentences: breadth (one database, small corpus), pass-3 depth, and the maintenance loop. Not sacrificed: the question, the scope, the logged search, critical reading of what you did read, notes, a matrix, an argued structure, and honest citation, which is to say, everything that makes a review a review. The emergency protocol is the full workflow's skeleton crew, not its impostor.

And one sentence for afterward, once the deadline passes and your pulse settles: the difference between this week and a calm one was never talent or time management mysticism; it was starting the system in week one instead of week eleven. Next project, start at Chapter 4 on day one, and this appendix becomes a souvenir.

A Note from the Author

Thank you for reading, and welcome to the conversation; it's better with you in it. Books like this one travel by word of mouth from students, so if the workflow rescued your semester, a short honest review where you bought it, or a loan to the labmate still drowning in tabs, genuinely matters.

The templates, the matrix starter, and the checklists are meant to be used, not admired: printable companions live at **gauravtiwari.org**, along with more of my writing for students. Corrections and workflow stories are welcome there too; future editions will owe their sharpest fixes to readers who wrote in.

Also by Gaurav Tiwari

How to Write Mathematics

Proofs, Papers, and Problem Sets for Students

The writing half of the craft: proofs graders can follow, full-marks problem sets, and your first paper, taught through before-and-after rewrites of real student work.

LaTeX for Theses and Dissertations

A Complete System for Writing, Formatting, and Submitting Your Thesis

The document half: from empty folder to accepted upload, with every university rule encoded once and every what-breaks moment handled.

LaTeX for Students

A Practical Guide to Beautiful Documents and Mathematical Typesetting

The on-ramp: documents, equations, figures, and templates, step by step, for students with no prior experience.

Other titles

- *Indian Polity: Complete Notes for UPSC Prelims & Mains*
- *Website and Blog Copywriting*
- *Marketing Essentials for Bloggers*
- *The Affiliate Bloggers System*
- *Ultimate Guide to Backlink Building Strategies*

All titles are available on Amazon. Find the full, current catalog at **gauravtiwari.org**.